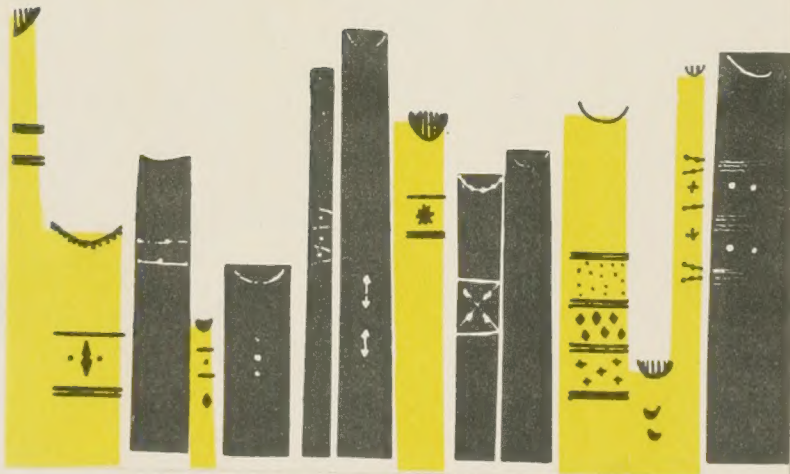


reference collection book



kansas city
public library
kansas city,
missouri



PUBLIC LIBRARY
SARASOTA CITY
FLA

From the collection of the

j f d
y z n m k
x o Preinger h
u v q g a
e s w p c
b t

San Francisco, California
2008

THE BELL SYSTEM

Technical Journal

DEVOTED TO THE SCIENTIFIC AND ENGINEERING
ASPECTS OF ELECTRICAL COMMUNICATION

ADVISORY BOARD

A. B. GOETZE

M. J. KELLY

E. J. McNEELY

EDITORIAL COMMITTEE

B. McMILLAN, *Chairman*

K. E. GOULD

S. E. BRILLHART

E. I. GREEN

A. J. BUSCH

R. K. HONAMAN

L. R. COOK

H. R. HUNTLEY

A. C. DICKIESON

F. R. LACK

R. L. DIETZOLD

J. R. PIERCE

EDITORIAL STAFF

W. D. BULLOCH, *Editor*

R. L. SHEPHERD, *Production Editor*

INDEX

VOLUME XXXVI

1957

AMERICAN TELEPHONE AND TELEGRAPH COMPANY
NEW YORK

KANSAS CITY, MO.
PUBLIC LIBRARY

FEB 5 1958

LIST OF ISSUES IN VOLUME XXXVI

No. 1	January.....	Pages	1-348
“ 2	March.....		349-592
“ 3	May.....		593-830
“ 4	July.....		831-1046
“ 5	September.....		1047-1318
“ 6	November.....		1319-1514

Technical Journal

DEVOTED TO THE SCIENTIFIC AND ENGINEERING
ASPECTS OF ELECTRICAL COMMUNICATION

VOLUME XXXVI

JANUARY 1957

KANSAS CITY, MO.
PUBLIC LIBRARY
FEB 13 1957
NUMBER 1

Transatlantic Communications — An Historical Resume 1
MERVIN J. KELLY AND SIR GORDON RADLEY

Transatlantic Telephone Cable System—Planning and Over-All Performance 7
E. T. MOTTRAM, R. J. HALSEY, J. W. EMLING AND R. G. GRIFFITH

System Design for the North Atlantic Link 29
H. A. LEWIS, R. S. TUCKER, G. H. LOVELL AND J. M. FRASER

Repeater Design for the North Atlantic Link 69
T. F. GLEICHMANN, A. H. LINCE, M. C. WOOLEY AND F. J. BRAGA

Repeater Production for the North Atlantic Link 103
H. A. LAMB AND W. W. HEFFNER

Power Feed Equipment for the North Atlantic Link 139
G. W. MESZAROS AND H. H. SPENCER

Electron Tubes for the Transatlantic Cable System 163
J. O. McNALLY, G. H. METSON, E. A. VEAZIE AND M. F. HOLMES

Cable Design and Manufacture for the Transatlantic Submarine Cable System 189
A. W. LEBERT, H. B. FISCHER AND M. C. BISKEBORN

System Design for the Newfoundland-Nova Scotia Link 217
R. J. HALSEY AND J. F. BAMPTON

Repeater Design for the Newfoundland-Nova Scotia Link 245
R. A. BROCKBANK, D. C. WALKER AND V. G. WELSBY

Power-Feed System for the Newfoundland-Nova Scotia Link 277
J. F. P. THOMAS AND R. KELLY

Route Selection and Cable Laying for the Transatlantic Cable System 293
J. S. JACK, CAPT. W. H. LEECH AND H. A. LEWIS

Bell System Technical Papers Not Published in This Journal 327

Recent Bell System Monographs 335

Contributors to This Issue 338

THE BELL SYSTEM TECHNICAL JOURNAL

ADVISORY BOARD

A. B. GOETZE, *President, Western Electric Company*

M. J. KELLY, *President, Bell Telephone Laboratories*

E. J. McNEELY, *Executive Vice President, American Telephone and Telegraph Company*

EDITORIAL COMMITTEE

B. McMILLAN, *Chairman*

S. E. BRILLHART

A. J. BUSCH

L. R. COOK

A. C. DICKIESON

R. L. DIETZOLD

K. E. GOULD

E. I. GREEN

R. K. HONAMAN

H. R. HUNTLEY

F. R. LACK

J. R. PIERCE

G. N. THAYER

EDITORIAL STAFF

J. D. TEBO, *Editor*

R. L. SHEPHERD, *Production Editor*

T. N. POPE, *Circulation Manager*

THE BELL SYSTEM TECHNICAL JOURNAL is published six times a year by the American Telephone and Telegraph Company, 195 Broadway, New York 7, N. Y. F. R. Kappel, President; S. Whitney Landon, Secretary; John J. Scanlon, Treasurer. Subscriptions are accepted at \$5.00 per year. Single copies \$1.25 each. Foreign postage is 65 cents per year or 11 cents per copy. Printed in U. S. A.

THE BELL SYSTEM TECHNICAL JOURNAL

VOLUME XXXVI

JANUARY 1957

NUMBER 1

Copyright 1957, American Telephone and Telegraph Company

Transatlantic Communications— An Historical Resume

By DR. MERVIN J. KELLY* and SIR GORDON RADLEY†

(Manuscript received July 30, 1956)

The papers that follow describe the design, manufacture and installation of the first transatlantic telephone cable system with all its component parts, including the connecting microwave radio-relay system in Nova Scotia. The purpose of this introduction is to set the scene in which this project was undertaken, and to discuss the technical contribution it has made to the development of world communications.

Electrical communication between the two sides of the North Atlantic started in 1866. In that year the laying of a telegraph cable between the British Isles and Newfoundland was successfully completed. Three previous attempts to establish transatlantic telegraph communication by submarine cable had failed. These failures are today seen to be the result of insufficient appreciation of the relation between the mechanical design of the cable and the stresses to which it is subjected as it is laid in the deep waters of the Atlantic. The making and laying of deep sea cables was a new art and designers had few experiments to guide them.

During the succeeding ninety years, submarine telegraph communication cables have been laid all over the world. Cable design has evolved from the simple structure of the first transatlantic telegraph cable — a

* Bell Telephone Laboratories. † British Post Office.

stranded copper conductor, insulated with gutta-percha and finished off with servings of jute yarn and soft armoring wires — to the relatively complex structure of the modern coaxial cable, strengthened by high tensile steel armoring for deep sea operation. The coaxial structure of the conducting path is necessary for the transmission of the wide frequency band width required for many telephone channels of communication. The optimum mechanical design of the structure for this first transoceanic telephone cable has been determined by many experiments in the laboratory and at sea. As a result, the cable engineer is confident that the risk of damage is exceedingly small even when the cable has to be laid and recovered under conditions which impose tensile loads approaching the breaking strength of the structure.

The great difference between the transatlantic telephone cable and all earlier transoceanic telegraph cables is, however, the inclusion of submerged repeaters as an integral part of the cable at equally spaced intervals and the use of two separate cables in the long intercontinental section to provide a separate transmission path for each direction. The repeaters make possible a very large increase in the frequency band width that can be transmitted. There are fifty-one of these submerged repeaters in each of the two cables connecting Clarenville in Newfoundland with Oban in Scotland. Each repeater provides 65 db of amplification at 164 kc, the highest transmitted frequency. The working frequency range of 144 kc will provide thirty-five telephone channels in each cable and one channel to be used for telegraph traffic between the United Kingdom and Canada. Each cable is a one-way traffic lane, all the “go” channels being in one cable and all the “return” in the other.

The design of the repeaters used in the North Atlantic is based on the use of electron tubes and other components, initially constructed or selected for reliability in service, supported by many years of research at Bell Telephone Laboratories. Nevertheless, the use of so many repeaters in one cable at the bottom of the ocean has been a bold step forward, well beyond anything that has been attempted hitherto. There are some 300 electron tubes and 6,000 other components in the submerged repeaters of the system. Many of the repeaters are at depths exceeding 2,000 fathoms ($2\frac{1}{4}$ miles) and recovery of the cable and replacement of a faulty repeater might well be a protracted and expensive operation. This has provided the incentive for a design that provides a new order of reliability and long life.

On the North Atlantic section of the route, the repeater elements are housed in flexible containers that can pass around the normal cable

laying gear without requiring the ship to be stopped each time a repeater is laid. The advantages of this flexible housing have been apparent during the laying operations of 1955 and 1956. They have made it possible to continue laying cable and repeaters under weather conditions which would have made it extremely difficult to handle rigid repeater housings with the methods at present available.

A single connecting cable has been used across Cabot Strait between Newfoundland and Nova Scotia. The sixteen repeaters in this section have been arranged electronically to give both-way amplification and the single cable provides "go" and "return" channels for sixty circuits. "Go" and "return" channels are disposed in separate frequency bands. The design is based closely on that used by the British Post Office in the North Sea. Use of a single cable for both-way transmission has many attractions, including that of flexibility in providing repeatered cable systems, but no means has yet been perfected of laying as part of a continuous operation the rigid repeater housings that are required because of the additional circuit elements. This is unimportant in relatively shallow water, but any operation that necessitates stopping the ship adds appreciably to the hazards of cable laying in very deep water.

The electron tubes used in the repeaters between Newfoundland and Scotland are relatively inefficient judged by present day standards. They have a mutual conductance of 1,000 micromhos. Proven reliability, lower mechanical failure probability and long life were the criteria that determined their choice. Electron tubes of much higher performance with a mutual conductance of 6,000 micromhos are used in the Newfoundland-Nova Scotia cable, and it is to be expected that long repeatered cable systems of the future will use electron tubes of similar performance. This will increase the amplification and enable a wider frequency band to be transmitted; thus assisting provision of a greater number of circuits. If every advantage is to be taken of the higher performance tubes, it will be necessary to duplicate (or parallel) the amplifier elements of each repeater, in the manner described in a later paper, in order to assure adequately long trouble-free performance. This has the disadvantage of requiring the use of a larger repeater housing.

During the three years that have elapsed since the announcement in December, 1953 by the American Telephone and Telegraph Company, the British Post Office, and the Canadian Overseas Telecommunication Corporation, of their intention to construct the first transatlantic telephone cable system, considerable progress has been made in the development and use of transistors. The low power drain and operating voltage required will make practicable a cable with many more sub-

merged repeaters than at present. This will make possible a further widening of the transmission band which could provide for more telephone circuits with accompanying decrease in cost per speech channel or the widened band could be utilized for television transmission. Much work, however, is yet to be done to mature the transistor art to the level of that of the thermionic electron tube and thus insure the constancy of characteristics and long trouble-free life that this transatlantic service demands.

The present transatlantic telephone cable whose technical properties are presented in the accompanying papers, however, gives promise of large reduction in costs of transoceanic communications on routes where the traffic justifies the provision of large traffic capacity repeatered cables. The thirty-six, four-kilocycle channels which each cable of the two-way system provides, are the equivalent of at least 864 telegraph channels. A modern telegraph cable of the same length without repeaters would provide only one channel of the same speed. The first transatlantic telegraph cable operated at a much slower speed, and transmitted only three words per minute. The greater capacity of future cables will reduce still further the cost of each communication circuit provided in them. Such considerations point to the economic attractiveness, where traffic potentials justify it, of providing broad band repeatered cables for all telephone, telegraph and teletypewriter service across ocean barriers.

The new transatlantic telephone cable supplements the service now provided by radio telephone between the European and North American Continents. It adds greatly to the present traffic handling capacity of this service. The first of these radio circuits was brought into operation between London and New York in 1927. As demands for service have grown, the number of circuits has been increased. We are, however, fast approaching a limit on further additions, as almost all possible frequency space has now been occupied. The submarine telephone cable has come therefore at an opportune time; further growth in traffic is not limited by traffic capacity.

Technical developments over the years by the British Post Office and Bell Telephone Laboratories have brought continuing improvement in the quality, continuity and reliability of the radio circuits. The use of high frequency transmission on a single side band with suppressed carrier and steerable receiving antenna are typical of these developments. Even so, the route, because of its location on the earth's surface, is particularly susceptible to ionospheric disturbances which produce quality deterioration and at times interrupt the service completely.

Cable transmission will be free of all such quality and continuity limitations. In fact, service of the quality and reliability of the long distance service in America and Western Europe is possible. This quality and continuity improvement may well accelerate the growth in transatlantic traffic.

The British Post Office and Bell Telephone Laboratories are continuing vigorous programs of research and development on submarine cable systems. Continuing technical advance can be anticipated. Broader transmission bands, lower cost systems and greater insurance of continuous, reliable and high quality services surely follow.

Transatlantic Telephone Cable System — Planning and Over-All Performance

By E. T. MOTTRAM,* R. J. HALSEY,† J. W. EMLING*
and R. G. GRIFFITH‡

(Manuscript received October 10, 1956)

The transatlantic telephone cable system was designed as a link connecting communication networks on the two sides of the Atlantic. The technical planning of the system and the objectives set up so that this role would be fulfilled, are the principal subjects of this paper. Typical performance characteristics illustrate the high degree with which the objectives have been realized. Optimum application of the experience of the British Post Office with rigid repeaters and the Bell System with flexible repeaters, together with close cooperation among three administrations, have played a large part in achieving the objectives.

INTRODUCTION

The transatlantic telephone cable system was planned primarily to connect London to New York and London to Montreal, and thus serve as an interconnection between continent-wide networks on the two sides of the Atlantic. Thus, the system has to be capable of serving as a link in wire circuits as long as 10,000 miles, connecting telephone instruments supplied by various administrations and used by peoples of many nations. This role as an intercontinental link has, therefore, been a controlling consideration in setting the basic objectives for the system.

The end sections of the system utilize facilities which are integral parts of the internal networks of the United States, Great Britain and Canada, but the essential new connecting links, extending between Oban, Scotland, and the United States-Canada border, and forming the greater part of the system, were built under an Agreement between the joint owners — the American Telephone and Telegraph Company and its subsidiary the Eastern Telephone and Telegraph Company (operating in Canada), the British Post Office, and the Canadian Overseas

* Bell Telephone Laboratories. † British Post Office. ‡ Canadian Overseas Telecommunication Corporation.

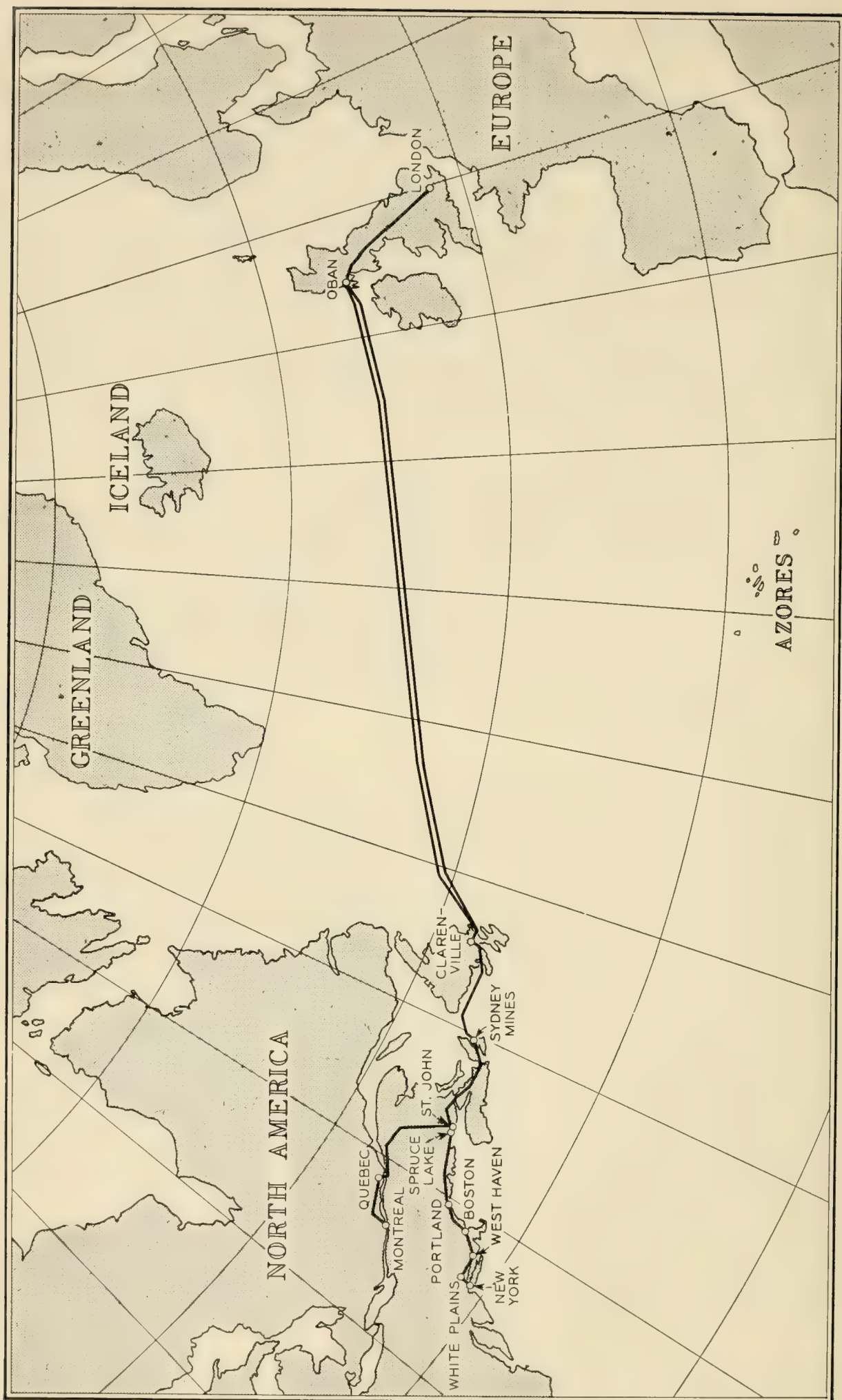


Fig. 1 — Route of the transatlantic submarine telephone cable system.

Telecommunication Corporation. It is thus the joint effort of three nations.

In planning the system, the main centres of interest were, naturally, the two submarine cable sections, Scotland to Newfoundland, and Newfoundland to Nova Scotia, each of which had to meet a unique combination of requirements imposed by water depth, cable length and transmitted bandwidth.

OVER-ALL VIEW OF THE SYSTEM

The transatlantic system provides 29 telephone circuits between London and New York, six telephone circuits between London and Montreal, and a single circuit split between London — New York and London — Montreal; this split circuit is available for telegraph and other narrow band uses. There are also 24 telephone circuits available for local service between Newfoundland and the Mainland of Canada, and there is considerable excess capacity over the radio-relay link that crosses the Maritime Provinces of Canada.

A map of the system is shown in Fig. 1; the facilities used, together with the approximate route distances are shown in Fig. 2. It will be seen that the over-all lengths of the London to New York and London to Montreal circuits are 4,078 and 4,157 statute miles respectively. Seven of the New York to London circuits are permanently extended to European Continental centres — Paris, Frankfurt (2), Amsterdam, Brussels, Copenhagen and Berne. The longest circuit is thus New York to Copenhagen, 4,948 miles.

Starting at London, which is the switching centre for United Kingdom and Continental points, 24-circuit carrier cables provide two alternative routes to Glasgow and thence to Oban by a new coaxial cable. Between London and Oban the two routes are fed in parallel at the sending ends, so a changeover can be effected at the receiving ends only. At a later date, an alternative route out of Oban will be provided by a new coaxial cable to Inverness.

From Oban a deep-sea submarine link connects to Clarenville, Newfoundland. This link is in fact two parallel submarine cables, one used for east-to-west transmission, the other for transmission in the reverse direction. Each cable is roughly 1,950 nautical-miles in length and lies at depths varying between a few hundred fathoms on the continental shelf and about 2,300 fathoms at the deepest point. Each cable incorporates 51 repeaters in flexible housings which compensate for the cable attenuation of about 3,200 db at the top frequency of 164 kc. These

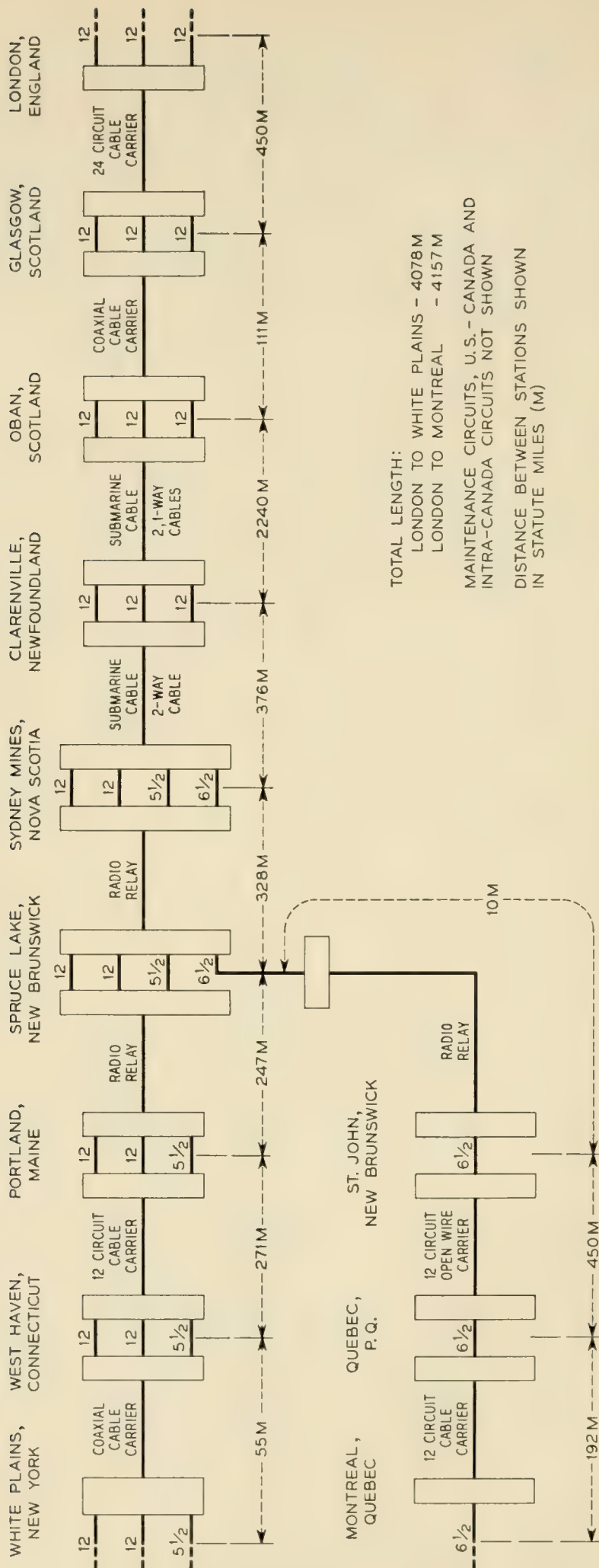


Fig. 2—Facilities used for the transatlantic submarine telephone cable system. Maintenance circuits, U.S.-Canada and intra-Canada circuits are not shown.

cables carry 36 telephone circuits plus maintenance circuits and establish the present maximum capacity of the transatlantic system.

At Clarendville, connection is made with Sydney Mines, Nova Scotia, by a second cable system which goes 63 statute-miles over land to Terrenceville, Newfoundland, and thence about 270 nautical-miles in coastal waters at a depth of about 250 fathoms. Although this system is partly on land, it is basically a submarine system in design, the two portions differing only in the protection of the cable. In this link, the two directions of transmission are carried by the same cable, a low-frequency band being used from west-to-east and a high-frequency band in the opposite direction. In addition to the necessary maintenance circuits, a total of 60 two-way circuits are provided, 36 being used for transatlantic service, and the remainder being available for service between Newfoundland and the Mainland. Sixteen two-way repeaters in rigid containers provide close to 1,000 db gain at this system's top frequency of 552 kc.

From Sydney Mines, transmission is by radio-relay to the United States-Canada border and thence to Portland, Maine; this system operates at about 4,000 mc and includes 17 intermediate stations. From Portland, standard 12-circuit carrier and coaxial cable facilities are used to connect with White Plains, New York, the American switching center 30 miles north of New York City, where connection is made to the Bell System network.

The Montreal circuits leave the radio-relay route at Spruce Lake, a relay station near the Border, from which point a short radio spur connects to St. John, New Brunswick, thence to Quebec on a 12-circuit open-wire carrier system and thence to Montreal on a 12-circuit cable carrier system.

BACKGROUND TO THE SUBMARINE CABLE SYSTEMS

The submarine cable sections have been built upon a long background of experience. Some of the cable laying and design techniques go back to the early telegraph cables of almost a century ago, and Lord Kelvin's analysis of the laying process is still the standard mathematical treatise on the subject. It is also interesting to note that the firm which provided most of the cable is a subsidiary of the organization that manufactured and laid the first successful transatlantic telegraph cable some 90 years ago.

In addition to the long experience in submarine telegraphy, the transatlantic system has drawn on over a quarter of a century of experience of telephone cable work in the British Post Office and the Bell System.

Experience in these two organizations has been quite different, but each in its own way has been invaluable in achieving today's system.

British Experience

In Great Britain, communication to the Continent dominated the early work in submarine telephony and led to systems providing relatively large numbers of circuits over short cables laid in shallow water. Early systems were un-repeatered, but the advantages of submerged repeaters were apparent. Experimental work, started in 1938, culminated in the first submerged repeater installation in an Anglesey-Isle of Man cable in 1943. Currently, there are many repeaters in the various shallow water cables radiating from the British Isles.

These repeaters, although of a size and mechanical structure well suited to shallow water applications, are not structurally suited to Atlantic depths. In 1948, the Post Office began to study deep water problems, and the first laying tests of a deep-water repeater housing were conducted in the Bay of Biscay in 1951. This housing was rigid, like the shallow-water ones, but smaller and double-ended so that the repeater was in line with the cable. Thus the rotation of the repeater, which accompanies the twisting and untwisting of the cable as tension is increased and decreased during the laying operation could be tolerated. The housing now used by the British Post Office is basically the same as this early deep-water design, although minor modifications have been made to improve the closure and water seals.

A serious study of transatlantic telephony was begun by the Post Office in 1950 when a committee was set up to report on future possibilities of repeatered cables. As a result, it was decided in 1952 to engineer a new telephone cable to Scandinavia, 300 nautical-miles in length, as a deep-water prototype, even though the requirements of depth, length, and channel capacity all could have been met by existing shallow-water designs.

All of the Post Office submarine systems are alike in that they use but a single cable, the go and return paths being carried by different frequency bands. The adoption of this plan was greatly influenced by the conditions under which the art developed. Because North Sea and Channel cables were highly subject to damage from fishing operations, it was desirable to limit the effects of such damage as much as possible. A single cable system is obviously preferable under these circumstances to a system using separate go and return cables which could be put out of service by damage to either cable. Since these systems were designed for shallow

water use, the additional container size required for two-way repeaters was of no great moment compared to the advantages of a single-cable system.

United States Experience

In the United States, the cable art developed under very different circumstances. There was, of course, need for communication to Cuba, Catalina, Nantucket and other off-shore locations, some of which involved conditions similar to those existing around the British Isles. The application of carrier to several of these cables occurred at an early date, but the repeater art was not directed at these shallow water applications.

For many years, telephone communication to Europe had been an important goal and some thirty-five years ago a specific proposal was made by the Bell System to the Post Office for a single, continuously-loaded, nonrepeatered cable to provide a single telephone circuit across the Atlantic.

This system was never built, partly because of the economic depression of the early thirties and partly because short-wave radio was able to meet current needs. Cable studies and experiments in the laboratory and field were continued, however, and largely influenced subsequent developments. It was at this time that the physical structure of the cable now used in the transatlantic system was worked out. It was also at this time that the harmful effects of physical irregularities in the cable were demonstrated. As cables are laid in deep water, high tensions are developed which unwrap the armor wires that normally spiral about the central structure. As tension changes during the laying process, twisting and untwisting occurs which is harmless if distributed along the cable. But obstructions in the cable which prevent rotation, or any other process such as starting and stopping of the ship which tends to localize twisting, are likely to cause kinking of cable and buckling of the conductors.

By 1932, electronic technology had advanced to a point where serious consideration could be given to a wideband system with numerous long-life repeaters laid on the bottom of the ocean and powered by current supplied over the cable from sources on shore.

The hazardous effects of obstructions in the cable, demonstrated in early laying tests, indicated that the chances of a successful deep-sea cable would be greatest if the repeaters were in small-diameter, flexible housings which could pass through laying gear without stopping the ship and without restricting the normal untwisting and twisting of the

cable. The structure ultimately evolved, consisting of two over-lapping layers of abutting steel pressure rings within a flexible waterproof container, was an important influence on the electrical design, since it placed severe limitations on size and placement of individual components.

Because these repeaters were to lie without failure for many years on the ocean bottom, it was necessary either to provide a minimum number of components of the utmost reliability, or to provide duplicate components to take over in case of failure. The size limitation favored the former approach. Similarly, the need for small size and minimum number of components militated against the use of two-way repeaters with their associated directional filters.

Out of these considerations grew the Bell System approach to solving the transatlantic problem by the use of two cables, each with built-in flexible amplifiers containing the minimum number of components of utmost reliability and a life objective of 20 years or better.

It was not until the end of World War II that such a system could be tried. At this time it was decided to install a pair of cables on the Key West-Havana route to evaluate the transatlantic design which had evolved in the prewar years. After further laying trials, this plan was completed in May, 1950, with the laying of two cables. Each of these had three built-in repeaters lying at depths up to 950 fathoms. These cables, each about 120 nautical-miles in length, carry 24 telephone circuits. They have now been in continuous service for over six years without repeater failure or evidence of deterioration.

EARLY TRANSATLANTIC TECHNICAL DECISIONS

Early in 1952, negotiations concerning a transatlantic cable were again opened between the American Telephone and Telegraph Company and the British Post Office. As indicated above, at that time each party had been laying plans for such a system. Thus it became necessary to evaluate the work on each side of the Atlantic to evolve the best technical solution.

To do this, a technical team from the Post Office visited Bell Telephone Laboratories in the fall of 1952 to examine developments in the United States. The work of the preceding 30 years was reviewed in detail with particular emphasis on the development and manufacture of the 1950 Key West-Havana cables. This was followed by a visit to the Post Office by a Bell Laboratories' team to review similar work in Great Britain. Again the review was comprehensive, covering shallow-water systems as well as plans for deep-water repeaters. Each visit was characterized by a frankness and complete openness of discussion that is perhaps unusual in international negotiations.

As is apparent from the previous discussion, it was found that the basic features of a deep-water design had been completed by the Bell System. Not only had many of the components been under laboratory test for many years, but a complete system had been operating for $2\frac{1}{2}$ years between Havana and Key West. To use a phrase coined at the time, the design had proven integrity.

Because of the years of proof and the conservative approach adopted to assure long life, the design was far from modern. The electron tubes, for example, had characteristics typical of tubes of the late 1930's, when, in fact, they were designed. Similarly, other components were essentially of prewar design.

The Post Office, on the other hand, had pioneered shallow-water repeaters and were pre-eminent in this field. Their deep-water designs were still evolving and had not yet been subjected to the same rigorous tests as the Bell System repeaters. This later evolution, however, made possible a much more modern design. The electron tubes, for example, had a mutual conductance of 6,000 micromhos as compared to about 1,000 in the Bell System repeater, and thus had a potentiality for much greater repeater bandwidths.

It was apparent from these reviews that only the American design was far enough advanced to assure service at an early date. It also appeared to have the integrity so essential to such a pioneering and costly effort as a transatlantic cable. On the other hand, the more modern Post Office design had many elements of potential value. If deep-water laying hazards could be overcome and proof of reliability established, it gave promise of greater flexibility and economy for future systems.

It was on these grounds that Dr. Mervin Kelly for the Bell System and Sir Gordon Radley for the Post Office jointly recommended that the Bell System design be used for the long length and great depths of the Atlantic crossing and the Post Office design be used for the Newfoundland-Nova Scotia link where the shallower water afforded less hazard and better observation of this potentially interesting design. The decision to use the Post Office design was subject to technical review after deep-sea laying tests and further experience with circuits and components. This review, made in June of 1954, confirmed the soundness of the original recommendation.

SYSTEM PLANNING

Planning of the individual systems began as soon as the technical decision just mentioned had been reached. By the time administrative agreements had been reached and the contract signed on November 27, 1953, both parties were ready to set up system objectives and an

over-all system plan. This work, too, was accomplished by a series of technical meetings held alternately in the United States and the United Kingdom, with additional meetings in Canada.

At the first of these meetings, a decision of far-reaching importance was made. It was agreed that each technical problem would be solved as it arose in so far as possible on the best engineering basis, putting aside all considerations of national pride. Adherence to this principle did much to forward the technical negotiations.

The initial joint meeting was also responsible for establishing most of the basic performance objectives of the system. The target date for opening of service, December 1, 1956, had been settled even earlier and was, in the event, bettered by nearly 10 weeks.

Service Objectives

A statement of the manner in which the system would be used and the services to be provided was a necessary preliminary to establishing performance objectives.

It was agreed that the system should be designed as a connecting link between the North American and European long distance networks. As such it should be capable of connecting any telephone in North America (ordinarily reached through the Bell System or Canadian long distance networks) with any telephone in the British Isles or any telephone normally reached from the British Isles through the European continental network. The system would be designed primarily for message telephone service but consideration would be given to the provision of other services such as VF carrier telegraph, program (music), and telephotograph as permitted by technical and contractual considerations. It was also agreed that the two submarine cable links should be so planned that it would be possible to utilize the full bandwidth in any desired manner in the future. Thus, for example, repeater test signals should be outside the main transmission band.

All elements in the submarine cable systems were to be planned for reliable service over a period of at least 20 years.

Transmission Objectives

The term "objective" was used advisedly in describing the aims of the system. It was agreed that such objectives were not ironclad requirements but rather desirable goals which it was believed practical to attain with the facilities proposed. Reasonable departure from these goals, however, would not be reason for major redesign.

Since the transatlantic circuits were to connect two extensive networks, the broad objective was to add as little loss and other forms of impairment as practical. To this end, they were to be designed essentially to the standards of international circuits as defined by the C.C.I.F.* and of circuits connecting main switching points in national networks, as for example, "Regional Centers" in the Bell System network and "Zone Centers" in the Post Office network.

The possibility of increasing the circuit capacity of the system by using channel spacings less than 4 kc was obvious. It was decided, however, to adopt, initially at least, the 4-kc spacing commonly used by long distance systems on both sides of the Atlantic. This would make possible the use of standard multiplexing arrangements, and it was believed that the number of circuits provided would be adequate for the first few years of operation. It would undoubtedly be desirable to increase the number of circuits in later years, but a decision on the method to be used was left until completion of exploratory work on several methods which promised capacity increases with less degradation than narrow-band operation.

The decision to use standard terminal equipment led naturally to acceptance of the principle that the 36 circuits across the Atlantic would be assembled as three 12-channel groups in the range 60–108 kc and the 60 circuits between Newfoundland and Nova Scotia as five 12-channel groups and thence as a supergroup in the range 312–552 kc. These are standard modulation stages in the multiplexing arrangements for broad-band carrier system on both sides of the Atlantic. Two of the 12-channel transatlantic groups would be connected to New York and the third would be split to provide $6\frac{1}{2}$ circuits to Montreal and $5\frac{1}{2}$ to New York in accordance with the Agreement.

To provide for program circuits, three eastbound and three westbound channels in each of the three transatlantic groups would be made available when required; equipment would be provided to replace either two or three 4-kc message telephone channels by a music channel. In order to avoid the agreed group pilot frequencies and to provide service to Montreal, it was agreed to utilize the frequency bands 68–76 kc and 64–76 kc in the 12-channel groups for this purpose. Terminals of British Post Office design would be used at all points for translation between program and carrier frequencies. The normal Bell System terminals could not be used since they occupy the frequency ranges 80–88 kc and

* The International Consultative Committee on Telephony (C.C.I.F.) bases its recommendations on a circuit 2,500 km (1,600 miles) in length, with implied pro rata increases for noise impairment.

76–88 kc which are not compatible with the split group arrangement or with the 84.08 kc end-to-end pilot.

Net Loss

The nominal 1,000-cycle net loss objective between London and New York for calls switched to other long distance trunks at each end (i.e., the via net loss) was set at 0.5 db. For calls terminating at either New York or London, the loss would be increased by switching a 3.5-db pad in London, as recommended by the C.C.I.F., and a 2-db pad at New York as standard in the Bell System. Thus a New York to London call would have a net loss of 6 db.

Variations from these nominal net losses owing to temperature effects, lack of perfect equalization and regulation, etc., are to be expected and a standard deviation of 1.5 db was set as the objective for such variations in the absence of trouble. The allocation of this variation to the various links is shown in Table I.

It is interesting to note that a smaller variation was allocated to the submarine links than to the over-land links. It was believed that the more stable environment on the ocean floor would make it possible to meet the rather small variation assigned to these links.

While these loss variations are consistent with normal long distance trunk objectives, they would not be satisfactory if compandors were found necessary to meet the noise objectives, and it was agreed that any of the links lying between such compandors would have to meet objectives half as large as those in Table I.

Frequency Characteristics

For telephone message circuits, the frequency characteristic recommended by the C.C.I.F., Fig. 3, was adopted with the expectation that it could be bettered by a factor of two, since channel equipments would be included at the circuit terminals only, as described later.

TABLE I — STANDARD DEVIATIONS OF NET LOSS OBJECTIVE

Link	Standard Deviation (db)
New York-Portland	0.75
Portland-Sydney Mines	0.75
Sydney Mines-Clareville	0.5
Clareville-Oban	0.5
Oban-London	0.75
Total (Assuming rms addition)	
New York-London }	1.5
Montreal-London }	

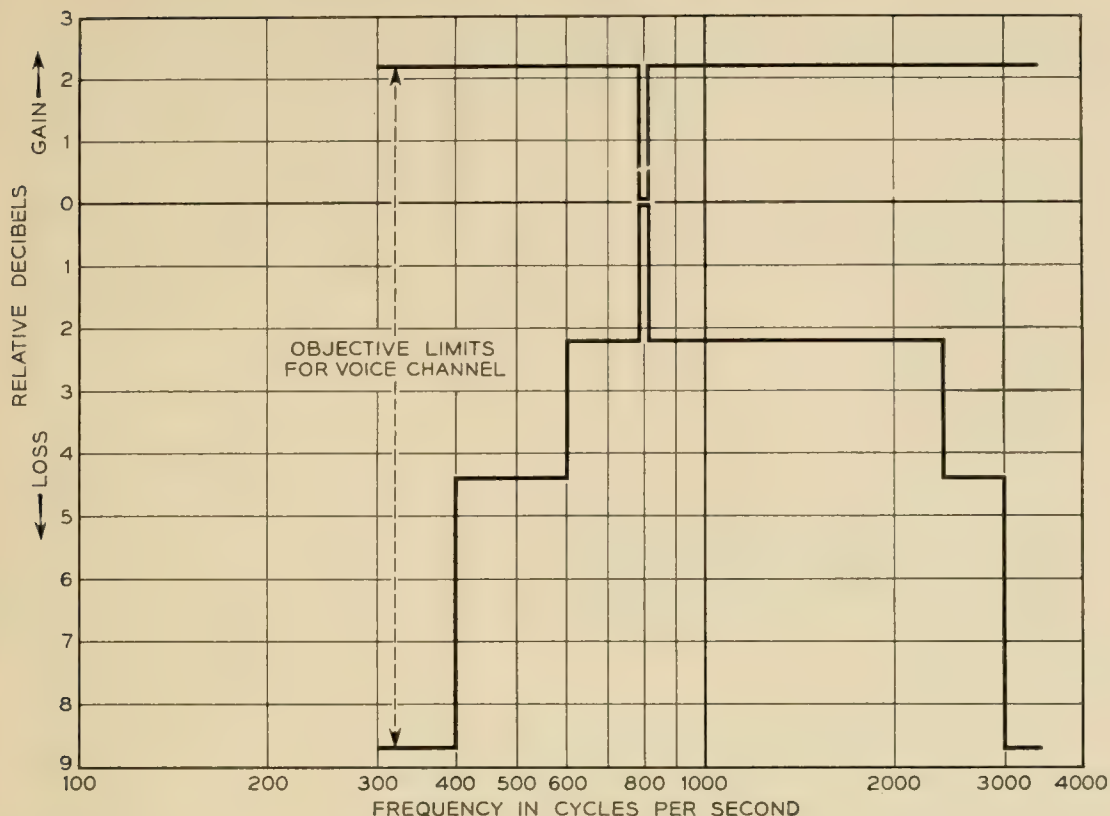


Fig. 3 — C.C.I.F. objectives for frequency characteristic of voice channel.

No specific objectives were agreed upon for the frequency characteristics of the 12-channel groups as such, but there was an expectation that ± 2 db could be achieved except for frequencies adjacent to the filters in the split group.

For program channels, the C.C.I.F. recommendations were also adopted in respect of the two-band (6.4 kc) and three-band (10 kc) arrangements. To meet the requirements of these channels and of telegraphy, an overall frequency stability objective of ± 2 cycles was adopted.

Noise and Crosstalk

Noise objectives were established to be reasonably consistent both with Bell System and C.C.I.F.* objectives for circuits of transatlantic length.

The objective for the rms circuit noise at a zero level point in the

* The methods specified by these two bodies for the assessment of circuit noise differ in three respects, the units employed, the frequency weighting employed, and the fraction of the busy hour for which the specified noise may occur. The meters concerned are the Bell System 2B noise meter (F1A weighting network) reading in dba and the C.C.I.F. psophometer (1951 weighting network) reading in millivolts across 600 ohms. The relationship between readings on the two meters is discussed in a later paper and it will suffice here to note that, for white noise $\text{dbm (CCIF)} = \text{dba (Bell)} - 84$.

TABLE II — RMS NOISE OBJECTIVES IN BUSY HOUR

Link	Approx. mileage	Noise dba
New York-Sydney Mines } Montreal-Sydney Mines }	1,000	31
Sydney Mines-Clareville.....	400	28
Clareville-Oban.....	2,000	36
Oban-London.....	500	28
Total		
New York-London } Montreal-London }		38

busy hour was agreed as 38 dba (i.e. -46 dbm or 3.9 mv). This was allocated between the various links as in Table II.

For the program channels, the agreed noise objective was -50 dbm as measured on a C.C.I.F. psophometer with a 1951 program weighting network.

Statistical data on probable speech levels and distributions at London and New York terminals were provided as a basis for repeater loading studies.

Early planning studies indicated that these objectives would probably be met on all, or nearly all channels without resort to compandors. If, as the system aged, the noise increased owing to increasing misalignment, the use of compandors would offer a means for reducing message circuit noise below the objectives.

The minimum equal-level crosstalk loss between any two telephone channels was set at 56 db for any source of potentially intelligible crosstalk. For channels used for VF telegraph, the equal-level crosstalk loss between go and return directions was set as a minimum of 40 db; for all program channels the minimum crosstalk attenuation would be 55 db.

Restrictions of Telegraph and Other Services

Since the system was being designed primarily for message telephone service, it was agreed that a channel used for any other service should not contribute more to the system rms or peak load than if this channel were used for message telephone, except by prior agreement between Post Office, Bell System and Canadian Overseas Telecommunication Corporation engineering representatives.

Signalling Objectives

In order to conserve frequency space, it was decided to transmit all calling and supervisory signals within the telephone channel bands and,

to avoid transmission degradation, it was agreed that the signaling power and duration would not amount to more than 9 milliwatt-seconds in the busy hour at a zero level point; this would not contribute unduly to the loading of the system.

It was agreed that, for initial operation, ringdown signaling would be employed, but the system design should be such as to permit the use of dialing at a later date.

Echo Suppressors

Echo control was considered essential, since the via net loss of the transatlantic circuits would be only 0.5 db, with a one-way transmission time of 35 milliseconds. Echo suppressors would be provided initially at New York and Montreal only, and arrangements made in London to cut out such suppressors as may be fitted there on Continental circuits, when these are used for extension of the transatlantic circuits. It was recognized, however, that other suppressors might be encountered in the more remote parts of Continental and United States extensions. The general problem of how best to arrange and operate echo suppressors on very long switched connections is one which remains for consideration later.

Maintenance and Operating Services

Telephone Speaker and Telegraph Printer Circuits

The need for telephone and telegraph circuits for maintenance and administration was recognized, and it was agreed to provide the following circuits on the submarine links at frequencies immediately outside the main transmission bands where inferior and somewhat uncertain characteristics might be expected (Fig. 4):

(a) A 4-kc band, possibly sub-standard in regard to noise, equipped with band splitting equipment (EB Banks) to provide two half-band-width telephone (speaker) circuits, and

(b) two frequency-modulated telegraph (printer) circuits.

These circuits would be extended over the land circuits to the terminal stations by standard arrangements as needed and would be used to provide the following facilities:

(I) An omnibus speaker circuit connecting the principal stations on the route, including Montreal.

(II) A speaker circuit for point-to-point communication between the principal stations — i.e., non-continuous.

(III) A direct printer circuit between London and White Plains.

(IV) An omnibus printer circuit as (I) above.

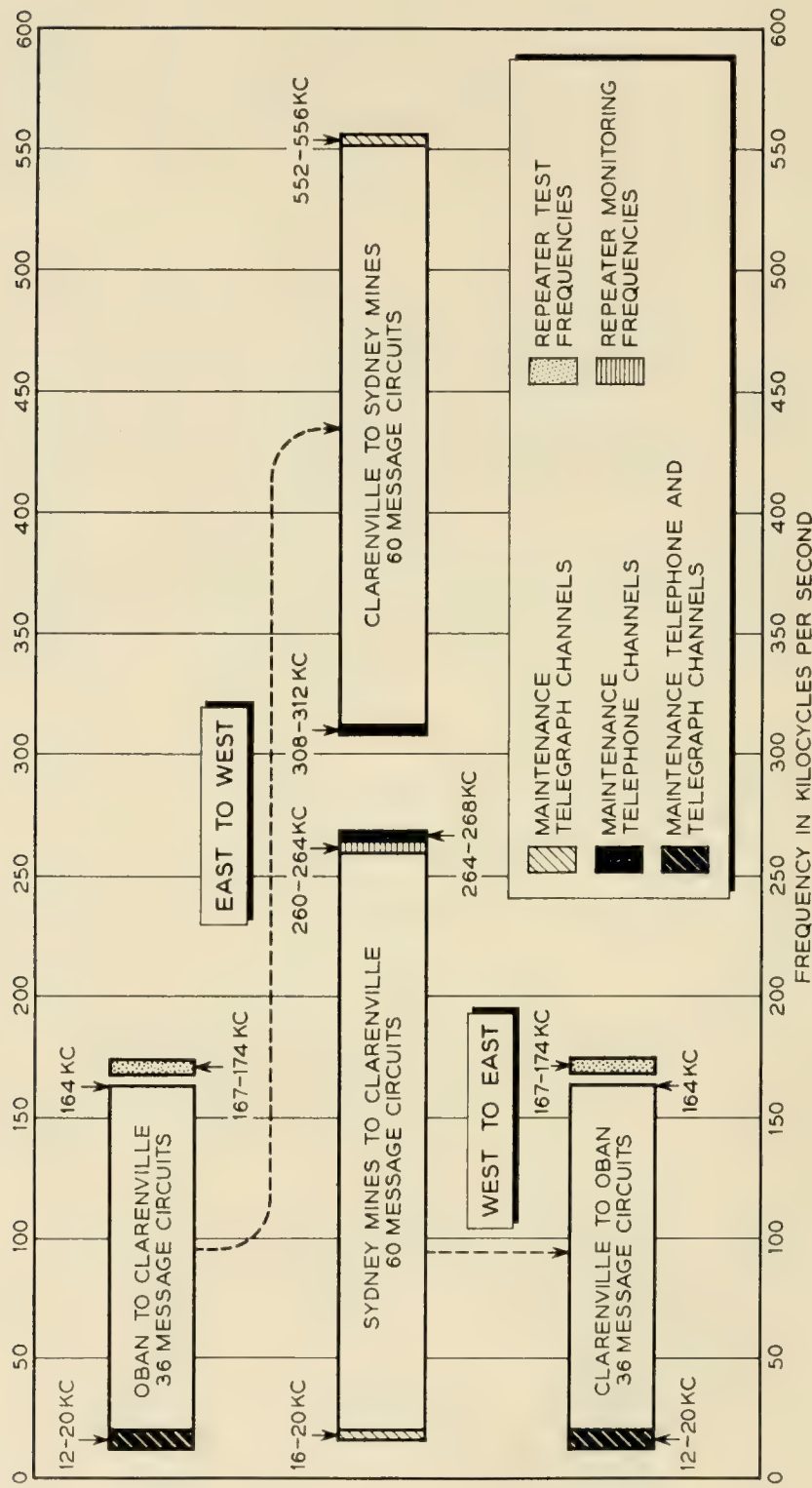


Fig. 4 — Frequency allocations in submarine cable links.

Repeater Test Frequencies

In each submarine cable link, test frequency bands were required for monitoring repeater performance; and these are indicated in Fig. 4.

Pilot Frequencies

It was agreed to provide pilot facilities throughout the route for line-up maintenance and regulation purposes. In addition to the usual pilots on the inland networks, there would be provided:

(a) a 92-kc pilot in each 12-channel group, continuous only in a particular section of the route and fitted with a recording voltmeter at the receiving end of that section, and

(b) an 84.08-kc overall pilot in each 12-channel group as recommended by the C.C.I.F. This would transmit continuously over the entire route and would be monitored and recorded at every main station.

Connections between Component Links

At the time that the objectives were being established, a far-reaching decision was made to employ channel equipment at London, New York, and Montreal only, and to adopt the frequency band 60–108 kc as the standard frequency for connecting the various parts of the over-all system. By adopting this band as standard for the transatlantic system, it also became possible to interconnect readily with land systems at each end.

This agreement also facilitated decisions on responsibility for design and manufacture of equipment. For example, it became logical to define each submarine system as the equipment between points where the 60–108-kc band appeared, i.e., the group connecting frames. Thus, these systems would include not only the cable, repeaters, and power supplies, but also the terminal gear to translate between 60–108 kc and line frequency of the submarine system. It also became logical to assign responsibility for manufacture of all of this equipment to the administration responsible for the specific system design, i.e., responsibility for the Oban-Clareville link to the Bell System and the Clareville-Sydney Mines link to the Post Office.

THE REALIZATION OF THE SYSTEM

With decisions reached on the system objectives and interconnecting arrangements, it became possible to lay out jointly a detailed over-all plan and for each administration to proceed with developing and engineering the links under its jurisdiction.

There was an understanding that there should be no deliberate attempt to make the characteristics of one link compensate for those of another, and so it would be incumbent on the administrations to produce the best possible group characteristic on each link.

The overall plan for the system, as finally developed, is shown in Fig. 2. Except for the necessity to split one of the three transatlantic groups in each direction to provide $6\frac{1}{2}$ circuits to Montreal and $5\frac{1}{2}$ to New York, which required specially designed crystal filters, no unusual circuit facilities were required.

Special equipment arrangements were called for at Sydney Mines and Clarendville to provide security for the Montreal-London circuits where they appeared in the same office with White Plains-London circuits. In these cases, a special locked room was constructed to house the equipment associated with the channel group containing the Canadian circuits.

The details of how the two all-important submarine cable links were designed and engineered to meet their individual objectives are given in companion papers. The efficiency and integrity of these two links are the highest that could be devised by engineers on both sides of the Atlantic.

Finally, each section of the connecting links was lined-up and tested individually before bringing them all together as an integrated system.

OVER-ALL PERFORMANCE OF THE SYSTEM

The system went into service on September 25, 1956, so soon after completion of some of the links that it was not possible to include all the final equalizers. Nevertheless, after completion of the initial overall line-up, the performance has been found to meet very closely the original objectives. The system went into service without the use of companders on any of the telephone circuits, but companders are included in the program equipment. At the time of writing, only the 2-channel program equipment is available for use.

Frequency Characteristics of 12-channel groups

Fig. 5 shows the frequency characteristic of one of the 12-channel groups, link by link and over-all, measured at group frequencies corresponding to 1,000 cycles on each channel. In both of the complete London-New York groups the deviation from flat transmission is within ± 1.5 db, and some further improvement is to be expected when the equalization is finalized. For the split group, the characteristics are similar except for the effect of the splitting filters.

Variation of Over-all Transmission Loss

The system has, of course, only been completed for a short time, but the indications so far are that the standard deviation of the transmission loss, as indicated by the 84.08-kc group pilots is well within the objective of 1.5 db. Alarms operate when the received pilot level deviates by ± 4 db and, so far, these alarms have not operated under working conditions.

Frequency Characteristics of Telephone Circuits

Fig. 6 shows the measured frequency characteristic of a typical circuit in the two directions of transmission as measured in the through and terminated conditions. Half the C.C.I.F. limits are met on most circuits.

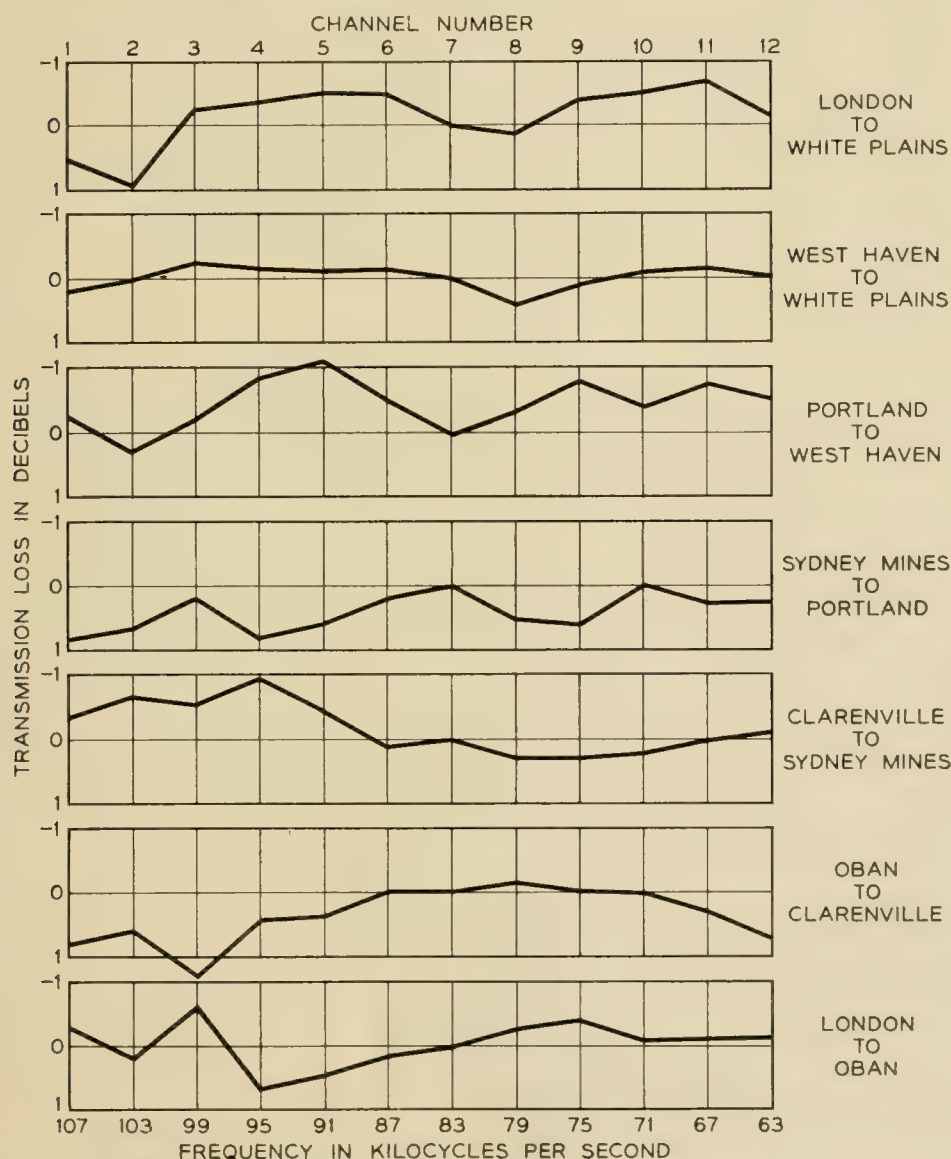


Fig. 5 — Frequency characteristic of typical London-New York channel group. (Measured at group frequencies corresponding to 1,000 cycles.)

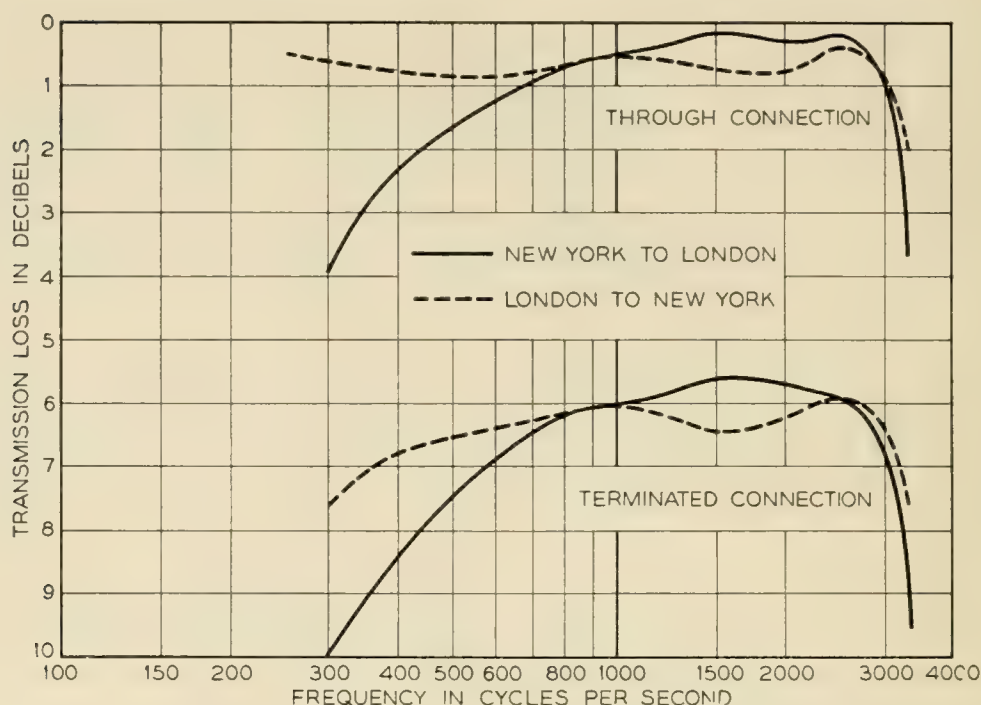


Fig. 6 — Frequency characteristic of typical telephone circuit.

Circuit Noise

The circuit noise, referred to a zero level point is as follows:

London-New York	Best 30 dba;	worst 36 dba
New York-London	Best 29 dba;	worst 41 dba
London-Montreal	Best 30 dba;	worst 33 dba
Montreal-London	Best 30 dba;	worst 31 dba

Two circuits at present exceed the objective of 38 dba in the New York-London direction only; the higher noise levels refer to the high frequency channels in the Oban-Clarendville cable. After additional data on the effect of cable temperature variations are accumulated, refinements will be made in the equalization and adjustment of levels on the Oban-Clarendville link. It is expected that the two worst channels can then be made to meet the objectives — still without the use of compandors.

Frequency Characteristics of Program Channels

Fig. 7 shows the measured frequency characteristic of a London-New York program channel; this is typical.

Telegraph Channels London-Montreal

In the Agreement it was envisaged that at least six 50-baud telegraph channels could be provided in each direction in the Canadian half circuit. In fact, eleven such channels have been provided using carriers spaced

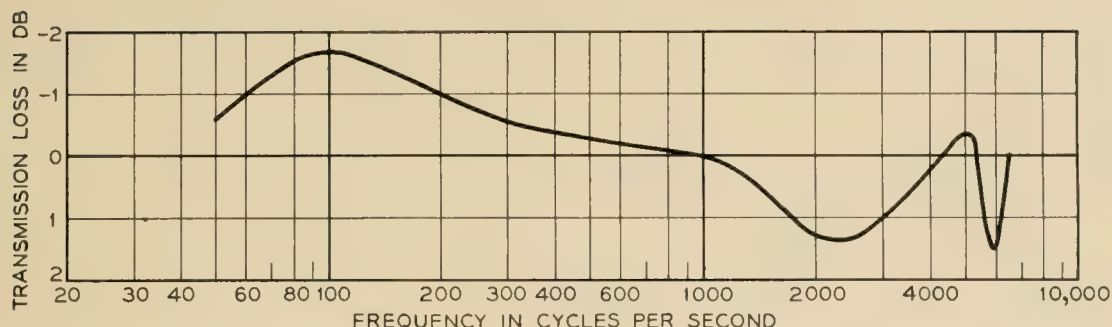


Fig. 7 — Frequency characteristic of typical program channel, London-New York.

at 120 cycles and frequency modulation. The telegraph distortion due to the cable system with start stop signals is about 4 per cent in every case, thus making the circuits suitable for switched connections, without regeneration, up to the same limits as inland systems.

Tests over the system indicate that the channel speed can be raised satisfactorily to 80 bauds on at least ten of the channels. By the adoption of synchronous working, it appears that time division multiplex systems can be operated on these ten channels to double their capacity at a later date.

CONCLUSION

The transatlantic cable system has presented unique problems in system planning and design. It has been necessary to design the system to connect the facilities of many countries and to provide for cable communication of unprecedented length. But the stringent design objectives necessary to meet these requirements have not been the only challenge to the designer. It has been necessary to meet these objectives with a system which for over 2,000 miles of its length could not be altered to the slightest extent once it had been placed on the ocean bottom. Except for the adjustments which can be made at the shore terminals of the submarine links it has not been permissible to make any of the multitude of small design changes, substitutions and adaptations which are so commonly required in new systems to achieve the design objectives.

The success achieved in meeting the original objectives is a measure of the realism of the early planning as well as the diligence with which the project was carried forward to completion and is a tribute to all who took part in planning, designing and building the system.

The accomplishment of getting into commercial service a working system with many complex links six weeks after the final splice was dropped overboard, and nearly ten weeks ahead of schedule, is a further tribute to the close cooperation of the technical people of three nations.

System Design for the North Atlantic Link

By H. A. LEWIS,* R. S. TUCKER* G. H. LOVELL* and
J. M. FRASER,*

(Manuscript received September 7, 1956)

The purpose of this paper is to examine the design and performance of the North Atlantic link, including consideration of factors governing the choice of features, a description of the operational design of the facility, and an outline of those measures available for future application in the event that faults or aging require corrective action.

DESCRIPTION OF LINK

That portion of the transatlantic system¹ which connects Newfoundland and Scotland consists of a physical 4-wire, repeatered, undersea link of Bell Telephone Laboratories design, with appropriate terminal and power feeding equipment in cable stations at Clarenville and at Oban.

The various elements comprising the link are shown in block form in Fig. 1. The termination points at each end of the link are the Group Distributing Frames (GDF), where the working channels are brought down in three groups to the nominal group frequency band 60–108 kc. At the west end, this point provides the interconnection between the North Atlantic and the Newfoundland-Nova Scotia links. At the east end, it is the common point between the North Atlantic link and the standard British toll plant over which the circuits are extended to London.

Two separate coaxial cables connect Clarenville with Oban, one handling east-to-west transmission, the other west-to-east. Each is about 1,940 nautical miles long. A total of 102 repeaters are installed in the two cables, at nominal intervals of 37.5 nautical miles. The cables also contain a number of simple undersea equalizers which are needed to bring system performance within the specified objectives.

The working spectrum of each cable extends from 20 to 164 kc, pro-

* Bell Telephone Laboratories.

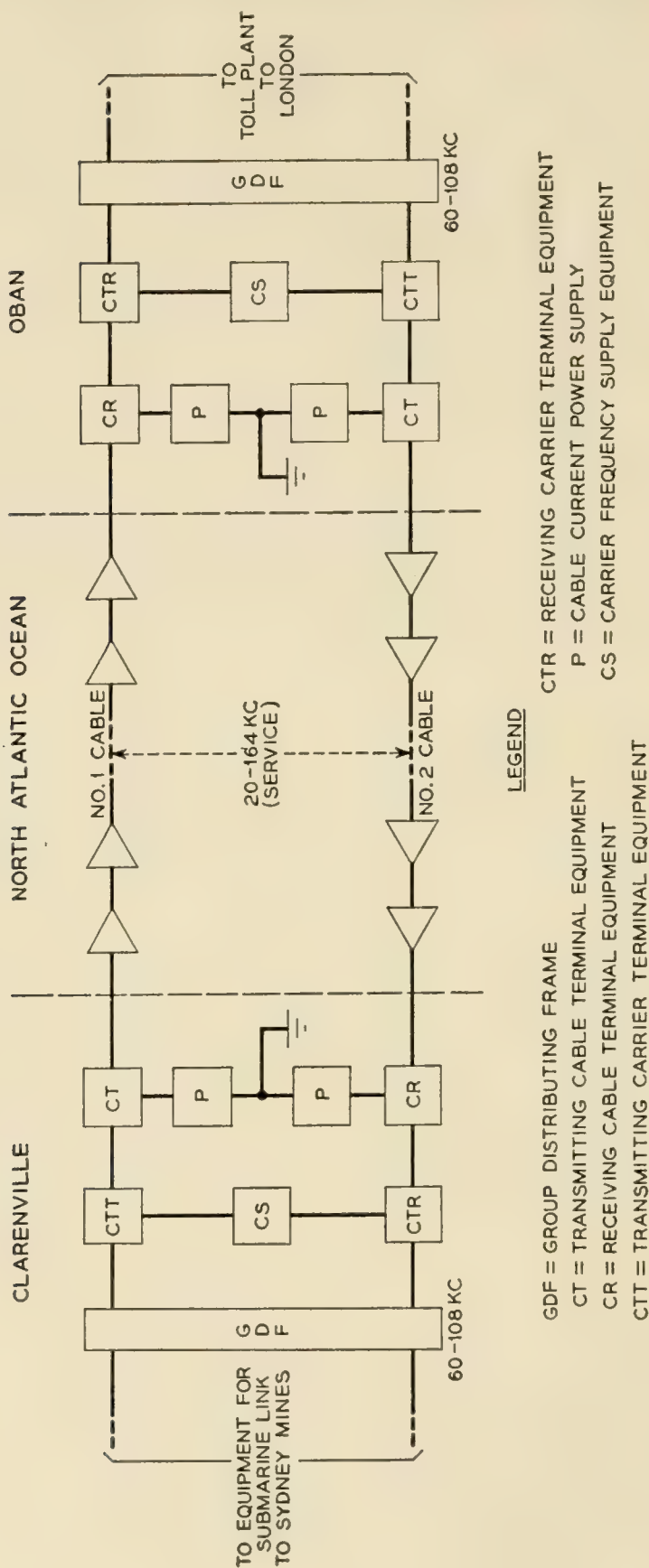


Fig. 1 — Schematic — North Atlantic link.

viding for 36 4-kc voice message channels. Below this band are assigned the telephone (speaker) and telegraph (printer) circuits needed for maintenance and administration of the facility. Above the working band, between 167 and 174 kc, are the crystal frequencies, which permit evaluation of the performance of each repeater individually from the shore stations.

The signal complex carried by the cables is derived in the carrier terminals from the signals on the individual voice frequency circuits by conventional frequency division techniques such as are employed in the Bell System types J, K and L broadband carrier systems manufactured by the Western Electric Company. (See Fig. 2.) These signals are applied to the cable through a transmitting amplifier which provides necessary gain and protects the undersea repeaters against harmful overloads. At the incoming end of each cable, a receiving amplifier provides gain and permits level adjustment.

Shore equalizers next to the transmitting and receiving amplifiers insert fixed shapes for cable length and level compensation. Adjustable units provide for equalization of the system against seasonal temperature changes on the ocean bed, and some aging.

The power equipment at each cable station includes (a) regular primary power with Diesel standby, (b) rotary machines for driving the cable current supplies, with battery standby, (c) battery plants for supplying the carrier terminal bays, and (d) last but by no means least in complexity, the cable current supplies themselves. These latter furnish regulated direct current to a series loop consisting of the central conductors of the two cables, with their repeaters. A power system ground is provided at the midpoint of the cable current supply at each cable station.

FACTORS AFFECTING SYSTEM DESIGN

General

A repeatered submarine cable system differs from the land-wire type of carrier system in two major respects. First, the cost of repairing a fault, and of the concurrent out-of-service time, is so great as to put an enormous premium on integrity of all the elements in the system and on proper safeguards in the system against shore-end induced faults. Second, once such a system is resting on the sea-bottom it is accessible for adjustment only at its ends. These two restrictions naturally had a profound influence on the design of the North Atlantic link.

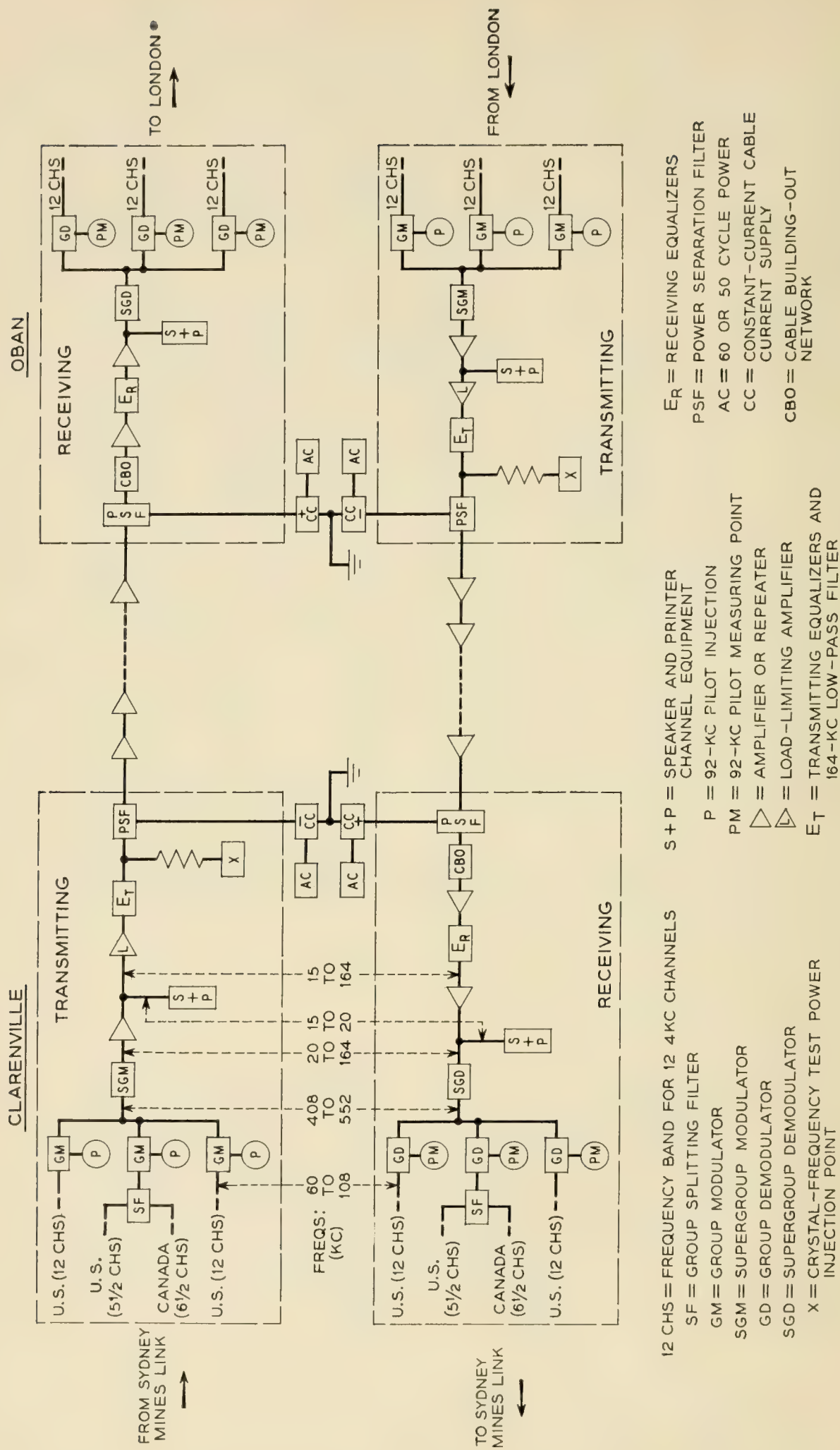


Fig. 2 — Block schematic of terminals — North Atlantic link.

Proven Integrity

Assurance of reliability dictated use of elements whose integrity had actually been proved by successful prior experience in the laboratory or in the Bell System. This resulted in adoption of a coaxial cable structure explored by the Laboratories soon after the war, field tested in the Bahamas area in 1948 and applied commercially on the USAF Missile Test Range project¹³ in 1953. It also resulted in specification of a basic flexible repeater design which had been under test in the Laboratories since before the war and had been in use on the Key West-Havana submarine cable telephone system⁴ since 1950.

Cable

The cable adopted was the largest which had actually been successfully laid in deep water, and its size afforded an important benefit in the form of reduced unit attenuation. The structure and characteristics of this cable are covered in detail in a companion paper.² Suffice it to point out here that by its adoption for the North Atlantic link, there were specified for the system designers (a) the attenuation characteristic of the transmission medium, (b) the influence on this of the pressure and temperature environments, (c) the unit contribution of the cable to the resistance of the power loop, and (d) the impedances faced by the repeaters.

Repeater

The adaptation³ of the basic Key West-Havana repeater design to the present project likewise presented the system designers with certain restrictions. Most important, the space and form of the long, tubular structure limited the size of the high voltage capacitors in the power separation filters, and thus their voltage rating, to the extent that this determined the maximum permissible number of repeaters. Likewise the performance of the repeater circuit was at least partially defined because of several factors. One was the effect of the physical shape on parasitic capacitances in the circuit, which in turn reacted on the feedback and hence on modulation performance, gain-bandwidth and the aging characteristic. Use of the Key West-Havana electron tube influenced the above factors and also tended to fix the input noise figure and load capacity.

System Design

In similar respects, the principle of proven integrity reacted into the broad consideration of system design. For instance, in a long system

having many repeaters in tandem, automatic gain control (gain regulation) in the repeaters provides an ideal method for minimizing the amount of the total system margin which must be allocated to environmental loss variations. However, this would have required adoption of elements of unproved integrity which might have increased the probability of system failure. So the more simple and reliable alternative was adopted — fixed gain repeaters with built-in system margins.

System Inaccessibility

The inaccessibility of the undersea system for periodic or seasonal adjustment and the decision to avoid automatic gain regulation were major factors in the allocation of system margins between undersea and shore locations. To avoid wasting such a valuable commodity as margin, required the most careful consideration of means of trimming the system during laying. Equalization for control of misalignment in the undersea link is a function of the match between cable attenuation and repeater gain. Generally speaking, these are fixed at the factory. Very small unit deviations from gain and loss objectives could well add to an impressive total in a 3,200-db system. Accordingly, it was necessary to plan for periodic adjustment of cable length during laying, and where necessary, insertion of simple mop-up undersea equalizers at the adjustment points.

DESIGN OF HIGH-FREQUENCY LINE

Terminology

The high-frequency line, as the term is used here, includes the cable, the undersea repeaters, the undersea equalizers, and the shore-station power separation filters, transmitting and line-frequency receiving amplifiers, and associated equalizers.

Repeater Spacing and System Bandwidth

As in most transmission media, the attenuation of the cable increases with increasing frequency. Hence the greater the bandwidth of the system, the greater the number of repeaters needed.

In this system, powered only from its ends, the maximum permissible number of repeaters and thus the repeater spacing is determined by a dc voltage limitation, as explained in a later section.

With the repeater spacing fixed, and the type of cable fixed, the required repeater gain versus frequency is known to the degree of accuracy that the cable attenuation is known. The frequency band which

can be utilized then depends on repeater design considerations, including gain-bandwidth limitations, signal power capacity and signal-to-noise requirements.

The early studies of this transatlantic system were based on "scaling up" the Havana-Key West system.⁴ In these studies, consideration was given to extending the band upward as far as possible by using companders⁵ on the top channels, thus lightening the signal-to-noise requirements on these channels by some 15 db provided they are restricted to message telephone service.

As the repeater design was worked out in detail, however, it became evident that a rather sharp upper frequency limit existed. This resulted from the parasitic capacitances imposed by the size and shape of the flexible repeater, the degree of precision required in matching repeater gain to cable loss in such a long system, and the feedback requirements as related to the requirement of at least 20 years' life.

These limitations resulted in the decision to develop a system with 36 channels of 4-kc carrier spacing, utilizing the frequency band from 20 to 164 kilocycles per second.

Signal-to-Noise Design

Scaling-up of Key West-Havana System

The length of the North Atlantic cable was to be about 16 times that of the Havana-Key West system. The number of channels was to be increased as much as practicable. The length increase entailed an increased power voltage to ground on the end repeaters. Increase in length and in number of channels entailed increased precision in control of variations in cable and repeaters. Work on these and other aspects was carried on concurrently, to determine the basic parameters of the extended system.

Increasing cable size decreases both the attenuation and the dc resistance, and in turn the voltage to ground on the end repeaters. It was soon decided that the largest cable size which could be safely adopted was the one used in the Bahamas tests. This has a center-conductor sea-bottom resistance of about 2.38 ohms per nautical mile. It consumes about 28 per cent of the total potential drop in cable plus repeaters.

Number of Repeaters

As indicated earlier, the factor which emerged as controlling the number of repeaters was the dc voltage to ground on the end repeaters. Considerations which entered into this were: voltage which blocking capacitors could safely withstand over a life of at least 20 years; volt-

age which other repeater elements such as connecting tapes between compartments could safely hold without danger of breakdown or corona noise; initial power potential and possible need for increasing dc cable current later in life to combat repeater aging; allowance for repair repeaters; and a reasonable allowance for increased power potential to offset adverse earth potential.

- Let R = dc resistance of center conductor (ohms/nautical mile)
 L = length of one cable* (nautical miles)
 E_{rep} = voltage drop across one repeater at current I
 I = ultimate (maximum) line current (amperes)
 N = ultimate number of submerged repeaters, in terms of equivalent regular repeaters
 n = allowance for repair repeaters, in terms of number of regular submerged repeaters using up same voltage drop.
 E_m = maximum voltage to ground at shore-end repeaters at end of life, and in absence of earth potential.
 S = spacing of working regular repeaters (nautical miles)

Then for the ultimate condition

$$2E_m = LIR - 2SIR + NE_{rep} \quad (1)$$

in which the term $2SIR$ accounts for the sum of the cable voltage drops on the two shore-end cable sections; the sum of their lengths is assumed, for simplicity, to equal $2S$. This equation also neglects a small allowance (less than 0.6 volt per mile) for the voltage drop in cable added to the system during repair operations. Also

$$S(N - n + 1) = L \quad (2)$$

because the repeater spacing is determined by the number of working regular repeaters. From (1) and (2),

$$2E_m = LIR + NE_{rep} - 2LIR/(N - n + 1) \quad (3)$$

The allowance n for repair repeaters was determined after studies of cable fault records of transoceanic telegraph cables, including average number of faults per year and proportion of faults occurring in shallow water. If the fault occurs in shallow water — as is true in most cases† — the net length of cable added to the system and the resulting attenuation increase, are small. Several shallow-water faults might be permis-

* The length of a cable is greater than the length of the route because of the need to pay out slack. The slack allowance, which averages 5 per cent in deep water on this route, helps to assure that the cable follows the contour of the bottom.

† Because of trawler activity, ship anchors and icebergs.

sible without adding a repair repeater. Therefore it was decided to let $n = 3$. Since the repair repeater is a 2-tube repeater while the regular repeater has 3 vacuum tubes, $n = 3$ corresponds to about 5 repair repeaters per cable.

To determine the maximum voltage E_m , it was necessary to consider blocking capacitors and earth potentials.

Based on laboratory life tests, the blocking capacitor developed has an estimated minimum life of 36 years at 2,000 volts. It is estimated that life varies inversely as about the fourth power of the voltage. The potential actually appearing on these capacitors is determined by the distance of the repeater from shore, the power potential applied to the system, and the magnitude and polarity of any earth potential.

Earth potential records on several Western Union submarine telegraph cables were examined. These covered a continuous period from 1938 to 1947, including the very severe magnetic storms of April, 1938, and March, 1940. It was judged reasonable to allow a margin of 400 volts (200 volts at each shore station) for magnetic storms during the final years of life of the system.

With an assumed maximum voltage of 2,300 volts on the end repeaters due to cable-current supply equipment, and 200 volts per end as allowance for the maximum opposing earth potential which the system would be permitted to offset without automatic reduction of the cable current, the voltage across the end repeater in late years of life (i.e., after line current had been increased to offset aging) would normally be 2,300 and would infrequently rise to 2,500. On the rare occasions where earth potential would rise above twice 200 volts, the cable current would be somewhat reduced and the transmission affected to a reasonably small extent. In the early years of life when the cable-current supply voltage would normally be about 2,000 volts, an opposing earth potential of twice 500 volts could be accommodated without affecting cable current; according to the telegraph cable records this would practically never occur. With conditions changing in this way over the years, the life of the blocking capacitors in the end repeaters was calculated to be satisfactory.

Accordingly E_m was established as 2,300 volts.

The system length L was estimated as about 1,985 nautical miles and the ultimate current I was estimated as 0.250 amperes, with a corresponding ultimate E_{rep} of about 62.8 volts.

Substituting in (3),

$$N = 55$$

$$N-n = 52 \text{ working repeaters}$$

$$S = 37.4 \text{ nautical miles.}$$

From a later estimate of $L = 1,955$ nautical miles, S calculates to be 36.9. The repeater design was based on this spacing in the deep sea temperature and pressure environment. Subsequently, after better knowledge had been obtained of the cable attenuation in deep water, the actual repeater spacing for the main part of the crossing was changed to about 37.4 nautical miles for the eastbound (No. 1) cable and 37.6 for the westbound. Only 51 repeaters were required in each cable.

Number of Channels

The number of channels which could be transmitted was determined by the upper and lower boundary frequencies. In this system, the bottom frequency was established at 20 kc, primarily because of the loss characteristics of the power separation filters.³

Preliminary studies were made to estimate the usable top frequency. For a system having a fixed number of repeaters, this frequency falls where the maximum permissible repeater gain equals the loss of a repeater section of the cable — which varies approximately as the square root of frequency. The repeater gain is the difference between the repeater input and output levels. The minimum permissible transmission level at the repeater input depends on the random noise (fluctuation noise) contributed by cable and repeaters, and on the specified requirement for random noise. The maximum permissible transmission level at the repeater output may depend on the modulation noise contributed by the repeaters below overload, or on overload from the peaks of the multi-channel signal complex. In this system, overload was found to be the controlling factor.

An important consideration was to provide enough feedback in the repeater so that at the end of 20 years the accumulated gain change (Mu-Beta effect) in all the repeaters would not cause the signal-to-noise performance to fall outside limits. The usable feedback voltage was scaled from the Havana-Key West design according to the relation that this voltage varies inversely as the $\frac{5}{3}$ power of the top frequency. Electron tube aging was estimated from laboratory life tests on Key West-Havana type tubes.

Concurrently, detailed theoretical and experimental studies were being conducted on the transmission design of a repeater suited to transatlantic use with the chosen type of cable, as discussed in a companion paper.³ Intimate acquaintance with the repeater limitations led to a decision in 1953 to develop a system with a working spectrum of 144 kc (36 channels) and a top frequency of 164 kc. Use of compandors would not increase this top frequency appreciably.

Computed Noise

When the repeater design was established, the remaining theoretical work on signal-to-noise consisted in refining the determination of the repeater output and input levels; computations of system noise, and comparison of this with the objectives to establish the margin available for variations; and determination of the necessary measures in manufacturing and cable laying so that these margins would not be exceeded by the deviations from ideal conditions which would occur. These deviations assumed great importance, because they tended to accumulate over the entire length of the system, and because many of them were unknown in magnitude before the system was actually laid.

The repeater levels and the resultant system noise were computed as follows:

It was recognized that the output levels of different repeaters would differ somewhat at any given frequency. The maximum allowable output level of the highest-level repeater was computed by the criterion that the instantaneous voltage at its output grid would be expected to reach the load-limit voltage very infrequently. This is the system load criterion established by Holbrook and Dixon.⁶ It premises that in the busy hour, the load-limit voltage (instantaneous peak value) should be reached 0.001 per cent of the time, or less. It is probable that the level could be raised 2 db higher than the one computed in this way without noticeable effect on intermodulation noise.

An important factor in the Holbrook-Dixon method is the talker volume distribution. Because of the special nature of this long circuit, a careful study was made of the expected United States talker volumes.

First, recent measurements of volumes on long-distance circuits were examined for the relation between talker volume and circuit length. They showed a small increase for the longer-distance circuits. This relation was extrapolated by a small amount to reach the 4000-mile value appropriate to the New York-London distance.

Second, an estimate was made of the probable trends in the Bell System plant in the next several years, which might affect the United States volumes on transatlantic cable calls.

The result of this was a "most probable U. S. volume distribution". This distribution, which had an average value of -12.5 vu at the zero level point, with a standard deviation of 5 db, agreed very well with one furnished by the Post Office and based on calls between London and the European continent. A further small allowance was then made for the contribution of signaling tones and system pilots which brought the resultant distribution to an average value of -12 vu at zero level of the

system (the level of the outgoing New York or London or Montreal switchboard), and a standard deviation of 5 db. It is approximately a normal-law distribution (expressed in vu), except that the very infrequent high volumes are reduced by load limiting in the inland circuits.

The other data needed for system load computations are the number of channels, and the "circuit activity", i.e., the per cent of time during the busy hour that the circuit is actually carrying voice in a given direction (eastbound or westbound). The circuit activity value used in designing United States long-distance multi-channel circuits is 25 per cent; for the transatlantic system, 30 per cent was used.

The peak value of the computed system multi-channel signal is the same as the peak value of a sine wave having an average power of +17.4 dbm at the zero transmission level point of the system.

This value, together with the measured sine-wave load capacity of the undersea repeater, determines the maximum permissible output transmission level of the repeater. The measured sine-wave load capacity is about +13.5 dbm at 164 kc. Hence the maximum permissible output transmission level for the 164-kc channel is $+13.5 - 17.4 = -3.9$ db.

If the relative output levels of the various submarine repeaters were precisely known, the highest-level repeater could have an output transmission level of -3.9 db at 164 kc. An allowance of about 2 db was made for uncertainty in knowledge of repeater levels, giving -6 db as the design value for the maximum repeater output level at 164 kc.*

The repeater has frequency shaping in the circuit between the grid of its output tube and the cable. The maximum repeater output at lower frequencies is smaller than at 164 kc, but the maximum voltage on the output grid, from an overload standpoint, is approximately constant over the 20 to 164 kc band.

The transmission level of the maximum level repeater output is thus determined, based on load considerations. Another factor which might limit this level is modulation noise. This was found to be less of a restriction than the load limitation, however.

With the output level determined, the random noise and the modulation noise for the system can be computed. The random noise computation is made on the assumption that all repeaters are at the same level, and then a correction is made for the estimated differences in levels of the various repeaters. The equation is

$$N_0 = N_{in} + G - TL + 10 \log n + d_r$$

* At the time of writing, the No. 1 cable (eastbound) is set up with a somewhat lower maximum output level than this; the safe increase in level will be determined later.

where N_0 = system random noise, in dba,* referred to zero transmission level

N_{in} = random noise per repeater in dba, referred to repeater input level

G = repeater gain in db

TL = transmission level of repeater output

n = number of undersea repeaters

d_r = db increase in noise due to differing output levels of the various repeaters as compared to the highest level repeater.

At the top frequency of 164 kc, $TL = -6$ db as seen above, $G = 60.7$ db, and N_{in} is about -55.5 dba, which corresponds to a noise figure of about 2 db. Hence for 52 repeaters,

$$N_0 = -55.5 + 60.7 - (-6) + 10 \log 52 + d_r = 28.4 + d_r$$

At lower frequencies, the noise power referred to repeater input is greater because part of the equalization loss is in the input circuit; the repeater gain is less, to match the lesser loss of a repeater section of cable; and the repeater output level is lower on account of the equalization loss from output grid to repeater output. This is shown in Table I.

TABLE I

Channel	N_{in} = Approx. Input Noise, dba	G = Approx. Repeater Gain, db	TL = Approx. Output Trans Level, db	N_0 = Resulting Noise dba
Top Freq.....	-55.5	60.7	-6	$28.4 + d_r$
Middle.....	-53	45	-11	$20.2 + d_r$
Bottom Freq.....	-46	22	-19	$12.2 + d_r$

In order to estimate d_r (which is a function of frequency) before the cable was laid, the factors were studied which might contribute to differences between the levels of the various repeaters. Estimates were made of probable total misalignment — by which is meant the level difference between highest-level and lowest-level repeaters. These values together with estimates of the resulting noise increases, are shown in Table II. This is based on the assumption that the repeater levels would be distributed approximately uniformly between highest and lowest. The position of a repeater along the cable route is not significant.

* “dba” is a term used for describing the interfering effect of noise on a speech channel. Readings of the 2B noise meter with F1A weighting may be converted to dba by adding 7 db. dba may be translated to dbm (unweighted) by noting that flat noise having a power of 1 milliwatt over a 3,000-cycle band equals approximately +82 dba, and that 1 milliwatt of 1,000-cycle single-frequency power equals +85 dba.

TABLE II

Channel	M = Estimated Misalignment	d_r = Resulting Noise Increase	N_0 = Resulting System Random Noise (Approx.)
Top.....	12 db	7.4 db	36 dba
Middle.....	10 db	6.0 db	26 dba
Bottom.....	6 db	3.4 db	16 dba

A study was made of the expected intermodulation noise. This noise is affected by repeater level differences, by talker volumes, and by circuit activity. It has a time distribution, the most attention being given to the root-mean-square modulation noise in the busy hour.

Modulation noise was computed in two ways: by the Bennett method,⁷ and by the Brockbank-Wass method.⁸ Results by the two methods are in fairly good agreement. The values computed by the Bennett method are shown in Table III.

Because noise powers, not dba, are additive, these values of modulation noise would contribute only a very small amount to the total noise in the deep-sea cable system, as previously stated.*

Other possible contributory sources of noise in the deep-sea cable system are noise in the terminals, noise picked up in the shore lead-in cables, and corona noise in the repeaters. The system has been designed so that the expected total of these is small. Hence the estimate of total deep-sea cable system noise, made before cable-laying, gave the values shown in Table IV, to the nearest db.

Misalignment Control

Objective

Since the noise objective for the deep sea cable link was set as 36 dba in the joint meetings of the British Post Office and Bell System Representatives¹, the above estimate of system noise led in turn to the objective that the system be manufactured and laid in such a way that the misalignment, including effects of seasonal temperature change, should be no greater in the top channel than the value assumed in the estimate, which was 12 db. Half of this was allotted to initial misalignment and half to effects of temperature change. Greater misalignments were permissible at lower frequencies.

* The following table for adding two noises expressed in dba shows the magnitudes of the increases:

db difference between larger and smaller noises	2	4	6	8	10	15	20
Resulting db to be added to larger noise to get sum of the two	2.1	1.4	1.0	0.6	0.4	0.1	0+

TABLE III

Channel	Weighted Noise in dba at Zero Transmission Level		
	Second Order Products	Third Order Products	Total
Top.....	8.2	8.5	11.3
Middle.....	0.2	3.8	5.4
Bottom.....	-1.8	2.5	3.9

TABLE IV

Channel	Estimated System RMS Noise at 0 db Transmission Level	
	dba	dbm (Psophometer†)
Top.....	36	-48
Middle.....	26	-58
Bottom.....	16	-68

† The psophometer is the C.C.I.F. circuit noise meter. On message telephone circuits, dbm (psophometer) + 84 = dba using C.C.I.F. 1951 weighting and Bell System F1A weighting.

Control of this misalignment required extensive consideration in the equalization design of the system.

Causes of Misalignment

The basic causes of misalignment can be grouped as follows: those producing unequal repeater levels when the system is first laid; those resulting from changes in cable loss produced by changes in sea-bottom temperature; and those from aging of the cable or repeaters.

There are a large number of possible causes of initial misalignment. While the cable and repeaters were manufactured within very close tolerances to their design objectives, they could not exactly meet those objectives. In addition, the cable loss as determined in the factory must be translated to the estimated loss on sea-bottom, and the length of each repeater section must be tailored in the factory so that its expected sea-bottom loss will best match the repeater gain. Possible sources of error in this process include: uncertainty in temperature of the cable when it is measured in the factory, and in its temperature on the sea-bottom; uncertainty in temperature and pressure coefficients of attenuation; changes in cable loss between factory and sea-bottom conditions, not accounted for by pressure and temperature coefficients.

These latter changes in cable loss were called "laying effect". The determination of the magnitude of laying effect, and its causes, are dis-

cussed in the paper on cable design.² Suffice it to say here that after measuring the loss of a length of cable in the factory and computing the sea-bottom loss, the computed result was a little greater than the actual sea-bottom loss. The difference, i.e., the laying effect, was approximately proportional to frequency, and was greater for deep-sea than for shallow-water conditions. Its existence was confirmed by precise measurements on trial lengths of cable, made early in 1955 in connection with cable-laying tests near Gibraltar.

Laying effect had been suspected from statistical analysis of less precise tests of repeater sections of cable generally similar to transatlantic cable, as measured in the factory and as laid in the vicinity of the Bahamas. However, the repeater design had of necessity been established before the Gibraltar test results were known. The consequence was a small systematic excess of repeater gain over computed cable loss, approximately proportional to frequency, and amounting at the top frequency to about 0.04 db per nautical mile on the average, for deep-sea conditions, and about 0.025 db per mile for shallow-water conditions. While this difference appears small, it would accumulate in a transatlantic crossing of some 1,600 miles of deep sea and 350 miles of relatively shallow sea, to about 75 db at the top frequency. This would be enough to render almost half of the channels useless if no remedial action were taken.

After the system is laid, changes in cable loss or repeater gain may occur. These may be caused by temperature effects and by aging of cable or repeaters.

Comprehensive studies* of the expected amounts of temperature change were made both before and after the 1955 laying. These gave an estimate of 0 to ± 1 degree F annual variation in sea-bottom temperature in the deep-sea part of the route (about 1,600 nautical miles) and perhaps $\pm 5^\circ$ F on the Continental shelves (about 330 nautical miles). Use of these figures leads to a ± 5 -db variation in system net loss at 164 kc from annual temperature changes. If the deep-sea bottom temperature did not change at all, the estimated net loss variation due to temperature would be about half of this.

Control of Misalignment

The first line of defense against variations leading to misalignment was the design and production of complementary repeaters and cable. This is discussed in companion papers.

* Factual data on deep sea bottom temperatures are elusive. Many of the existing data were acquired by unknown methods under unspecified circumstances, using apparatus of unstated accuracies. Statistical analysis of selected portions of the data leads to the quoted estimates.

Signal-to-noise changes from undersea temperature variations were minimized by providing adjustable temperature equalizers at both the transmitting and the receiving terminals, and devising a suitable method of choosing when and how to adjust them.

Partial compensation for laying effect was carried out at the cable factory by slightly lengthening the individual repeater sections. In the 1955 cable, where the factory compensation was based on early data, the increased loss compensated for about two-fifths of the laying effect at the top frequency. In the 1956 cable, which had the benefit of the 1955 transatlantic experience, the compensation was increased to nearly twice this amount. Since the loss of the added cable is approximately proportional to the square root of frequency and the laying effect is approximately directly proportional to frequency, the proportion of laying effect compensated varied with frequency, and a residual remained which had a loss deficiency rising sharply in the upper part of the transmitted band.

The remainder of the laying effect, which was not compensated for in the factory, as well as other initial variations, were largely compensated by measures taken during cable laying at intervals of 150 to 200 miles. The whole length of the cable was divided into eleven "ocean blocks", each either four or five repeater sections long. At each junction between successive ocean blocks, means were provided to compensate approximately for the excess gain or loss which had accumulated up to that point.

These means were twofold. The first means was adjustment of the length of the repeater section containing the junction. For this purpose, the beginning of each block, except the first, was manufactured with a small excess length of cable. This was to be cut to the desired length as determined by shipboard transmission measurements.

The second means was a set of undersea equalizers. These equalizers were fixed series networks, encased in housings similar to the repeater housings but shorter.

Before the nature of the laying effect was known, it had been planned to have equalizers with perhaps several shapes of loss versus frequency; but to combat laying effect, nearly all were finally made with a loss curve sloping sharply upward at the higher frequencies, as shown in Fig. 3.* Because the equalizer components had to be manufactured many months in advance and then sealed into the housings, last minute designs of undersea equalizers were not practical. The proven-integrity principle prevented use of adjustable units. Because the equalizers were series-

* This characteristic was based on a statistical analysis of the data on some similar cable laid for the Air Force project.

type, each block junction was located at approximately mid-repeater section to minimize the effects of reflections between equalizer and cable impedances.

The actual adjustments at block junctions were determined by a series of transmission measurements during laying, as described in the companion paper on cable laying.⁹ Six equalizers were used at block junctions in the 1955 (No. 1) cable, and eight in the 1956 (No. 2) cable.

The result of all the precautions taken to control initial misalignment was, in the 1955 cable, to hold the initial level difference between highest and lowest-level repeaters to about 6 db near the top frequency and to values between 4 and 9 db at lower frequencies, the 9 db value occurring in the range 50 to 70 kc where there is noise margin. In the 1956 cable, the level difference was about 4 db at the top frequency, and from 2 to 7 db elsewhere in the band.

Shore Equalization

The equalization to be provided at the transmitting and receiving ends of the North Atlantic link had these primary functions:

1. For signal to noise reasons, to provide a signal level approximately flat with frequency on the grid of the third tube of the undersea repeaters.
2. To equalize the system so that the received signal level is approximately constant over the transmitted band.
3. To keep the system net loss flat, regardless of temperature variations in the ocean. (A change of 1° F in sea-bottom temperature would cause the cable loss to change by 2.8 db at 160 kc and less than this at lower frequencies. The amplifiers are relatively unaffected by small temperature changes.)
4. To provide overload protection for the highest level repeater.
5. To incorporate some adjustment against possible cable aging.

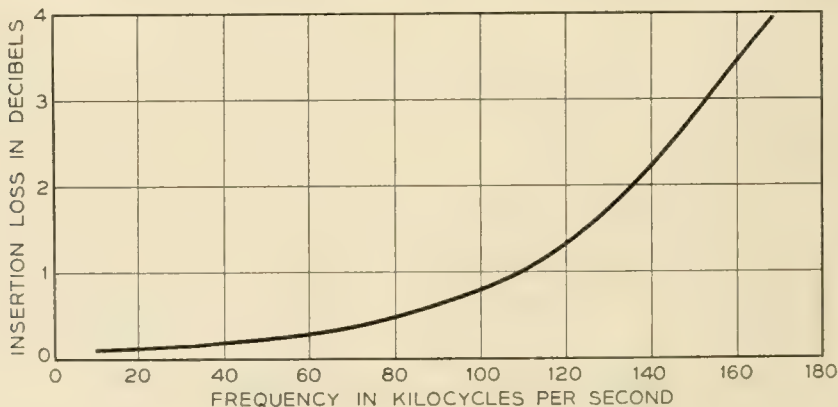


Fig. 3 — Undersea equalizer — loss-frequency characteristic.

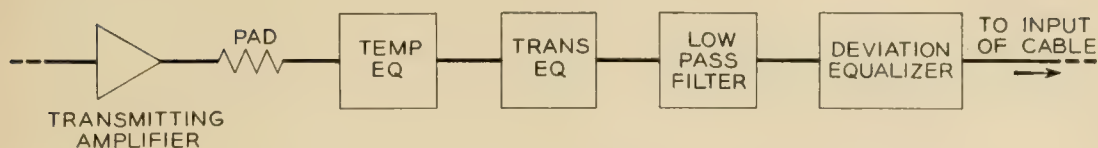


Fig. 4 — Transmitting equalizers — block schematic.

A portion of the transmitting terminal is shown in block schematic form in Fig. 4. All of the equalizers are constant resistance, 135 ohm unbalanced bridged-T structures. Their functions are as follows:

Transmitting Temperature Equalizer — It was estimated that the cable loss change from temperature variations would not exceed ± 5 db at the top of the transmitted band. The change in cable loss versus frequency due to temperature variations is approximately the same as if the cable were made slightly longer or shorter, and so the temperature equalizer was designed to match cable shape. Temperature equalizers are used both in the transmitting and receiving terminals to minimize the signal-to-noise degradation caused by temperature misalignment. Each equalizer provides a range of about ± 5 db at the top of the band, the loss being adjustable in steps of 0.5 db by means of keys. See Fig. 5.

This range might appear to be more than necessary, but it must be borne in mind that the sea bottom temperature data left much to be desired in precise knowledge of both average and range.

Transmitting Equalizer — The transmitting equalizer was so designed that with the temperature equalizer set at mid-range and with repeaters that match the cable loss, the signal level on the grid of the output tube of the first repeater would be flat with frequency. The loss-frequency characteristic is shown on Fig. 6.

L.P. Filter — This filter transmits signals up to 164 kc and suppresses by 25 db or more, signals in the frequency range from 165 to 175 kc. Its purpose is to prevent an accidentally applied test tone from overloading the deep sea repeaters should such a tone coincide with the resonance

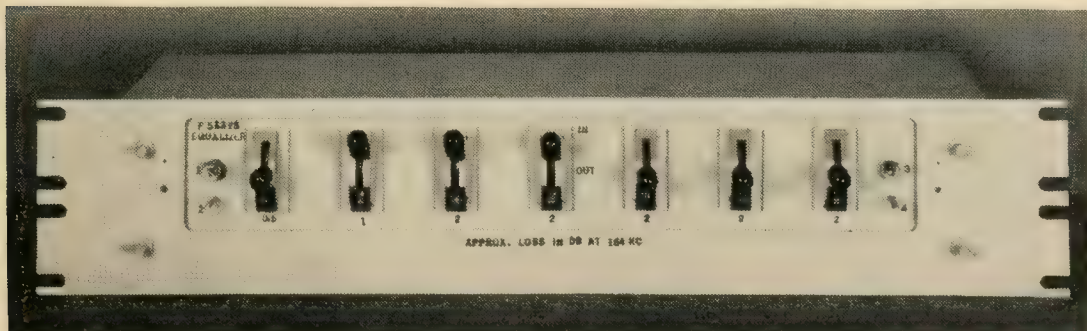


Fig. 5 — Transmitting temperature equalizer.

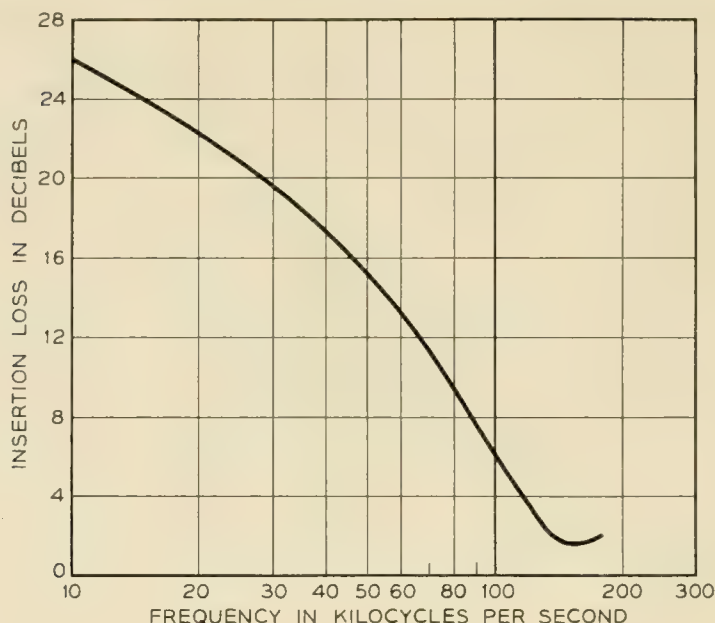


Fig. 6 — Transmitting equalizer — loss-frequency characteristic.

of the crystal used in a repeater for performance checking. The gain of a repeater to a signal applied at the resonance frequency of its crystal is about 25 db greater than at 164 kc.

Deviation Equalizer — The transmitting deviation equalizer is used to protect the highest level repeater from overload. The design is based on data obtained during the laying indicating which repeater was at the highest level at the various frequencies. The characteristic of this equalizer is shown on Fig. 7.

A portion of the receiving terminal is shown in block schematic form on Fig. 8. The generalized function of the equalizers is to make the signal level flat versus frequency at the input to the receiving amplifier No. 2. The specific functions of the various equalizers are as follows:

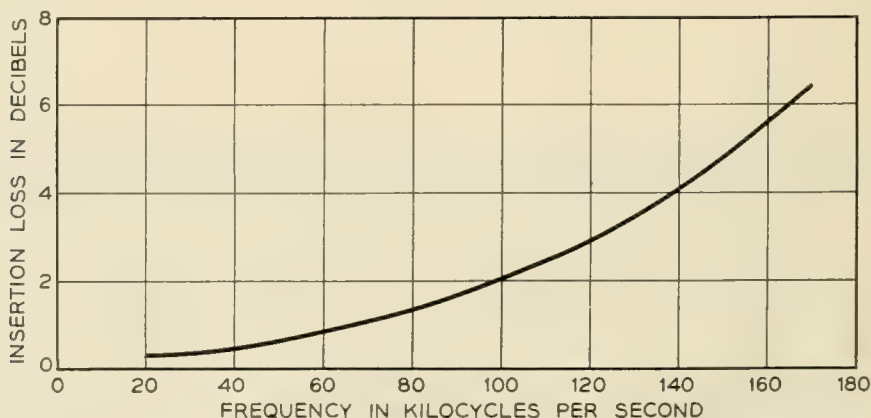


Fig. 7 — Transmitting deviation equalizer — loss-frequency characteristic.

Cable Length Equalizer — In planning the system, an estimate was made of how far the last repeater (receiving end) would be from the shore. Considerations of interference and level dictated a maximum distance not exceeding approximately 32 miles. For protection against wave action, trawlers and similar hazards, the repeater should be no closer than about five miles. To take care of this variation, a cable length equalizer was designed that is capable of simulating the loss of 10 miles of cable, adjustable in 0.5 db steps at the top of the frequency band. Two of these could be used if needed. Once the system is laid, this equalizer should require no further adjustment unless it is used to take care of cable aging or a cable repair near shore.

Receiving Fixed Equalizer — This is the mop-up equalizer for the system. A final receiving equalizer, Fig. 9, has been constructed for the first crossing. Another, tailored to the No. 2 cable, will be designed on the basis of data taken after completion.

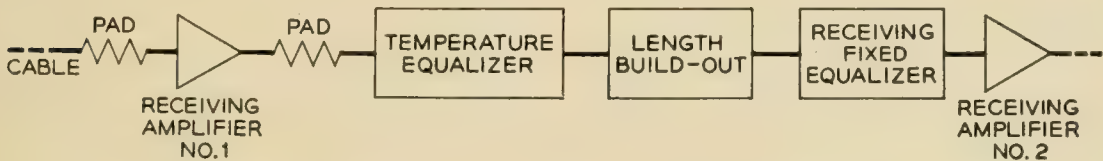


Fig. 8 — Receiving equalizers — block schematic.

Receiving Temperature Equalizer — The receiving temperature equalizer is identical with the transmitting temperature unit.

OPERATIONAL DESIGN

General

The operational design of a transmission system considers the supplementary facilities which are needed for operation of the main transmission facility, for its supervision and for its maintenance. In the present instance, these facilities include the cable station power plants for driving the carrier terminals and high frequency line, the carrier terminals themselves and their associated carrier supply bays, the telephone and telegraph (speaker and printer) equipments needed for maintenance, supervision and administration of the overall facility, the pilots, protection devices and alarms and the maintenance and fault locating equipment.

Power Supplies

With the exception of the plants for cable current supply, the power plants at Clarendville and Oban are relatively conventional, and follow telephone office techniques which are standard for the telephone adminis-

tration of the particular side of the Atlantic on which they are located. They will not be discussed here, although it might be well to point out that the equipment supply for the Bell equipment is direct current, obtained from floated storage batteries, while the supply for the Post Office equipment is alternating current from rotary machines driven normally from the station ac supply, with storage battery back-up.

The cable current supplies for the North Atlantic link¹⁰ are complex and highly special. It is pertinent to discuss here briefly the requirements which beget this complexity. To assist in this, an elementary schematic of the cable power loop is shown in Fig. 10.

The following main requirements governed the design of these plants:

1. Constant maintenance of uniform and known current in the power loop.
2. Protection of HF line against faults induced by failures in power bays.
3. Protection of HF line against damage from power voltage surges caused by faults in the line itself.

Maintenance of constant and known operating current in the line is very important from the standpoint of system life because of the dependence of the rate of electron tube aging on the power dissipated in the heaters.¹¹

The principal factor which tends to cause variations in the line current is the earth potential which may appear between the terminals of the system. This potential may be of varying magnitude and of either po-

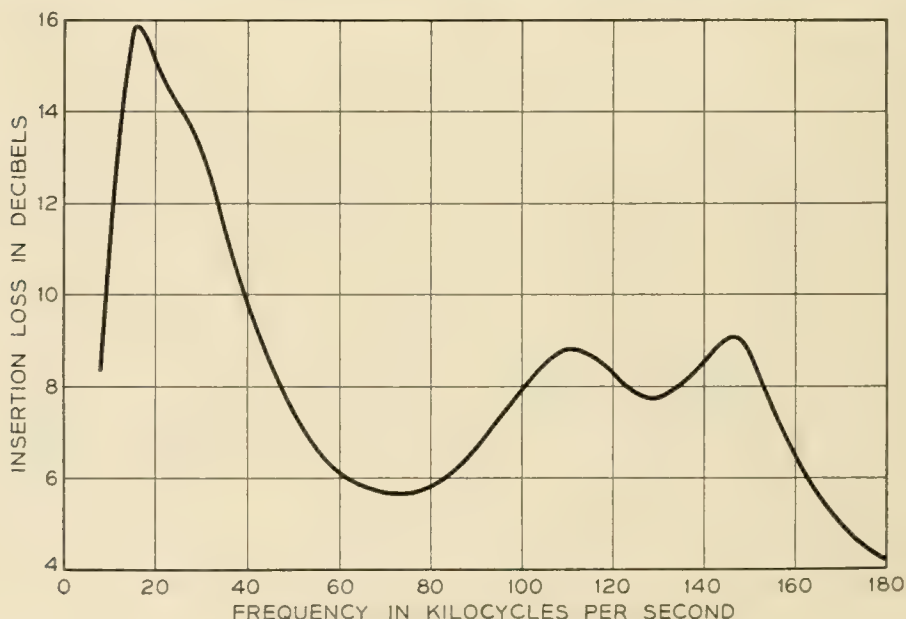


Fig. 9 — Receiving fixed equalizer — loss-frequency characteristic.

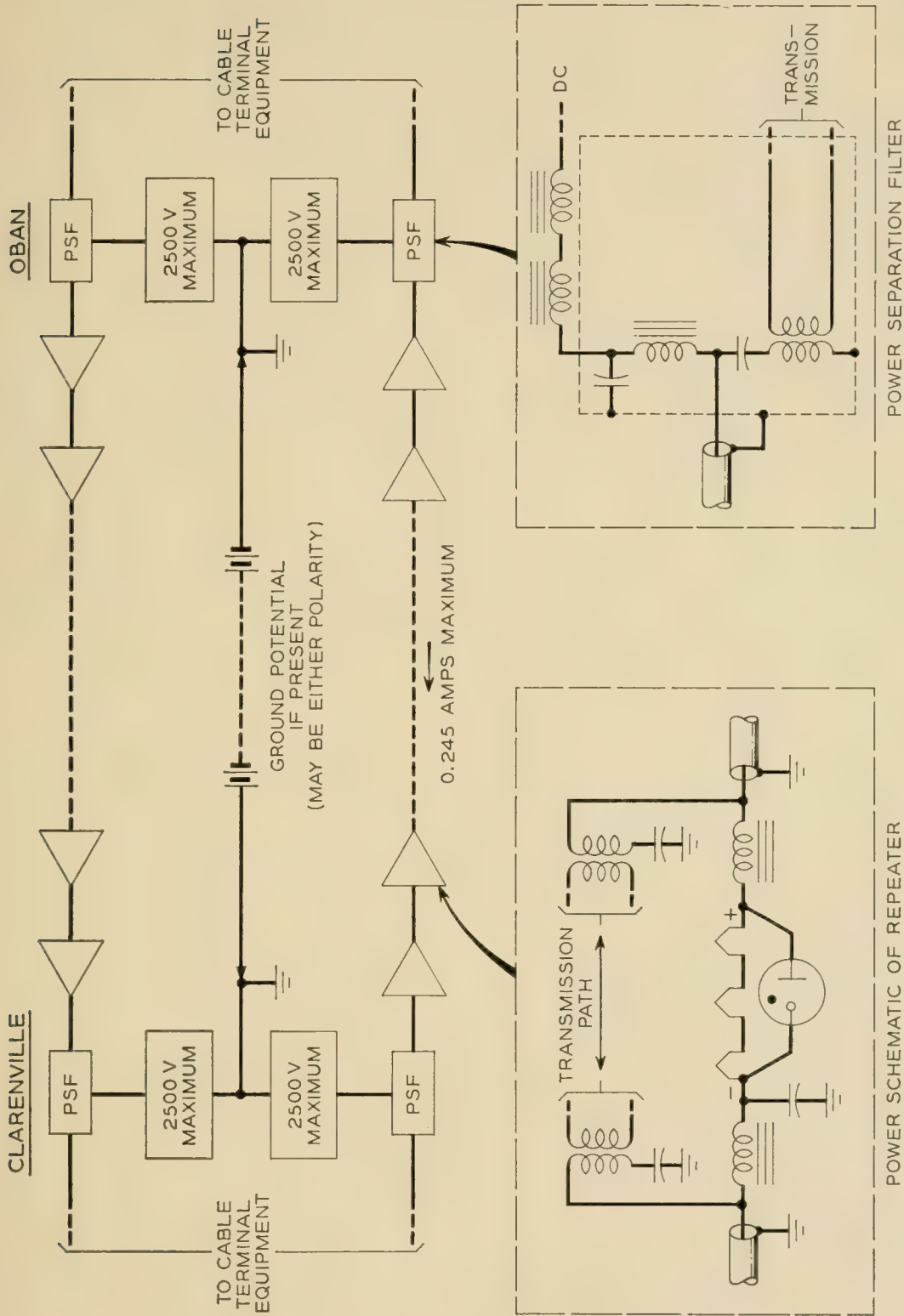


Fig. 10 — Cable power loop — elementary schematic.

larity. Consequently, it may either aid or oppose the driving potential applied to a particular cable. The power circuit regulation is such that the presence of an aiding potential will be completely compensated, while opposing potentials will be compensated only to the degree possible with the maximum of 2,500 volts applied by the terminal plants. Beyond this point, the line current will be allowed to drop, so as to avoid excessive potential across the power filter capacitors in the shore end repeaters.

Line faults caused by failures or misoperation of the power supplies must be carefully guarded against. Such failures might result in surges on the line, or excessive voltage or current. Protection against surges takes the form of retardation coils in the terminal power separation filters, which limit the rate of current change to tolerable values.

Failure of the capacitors in these filters could also create dangerous surges. The protection here takes the form of a large voltage design margin.

Faults in the HF line of importance from the power feed standpoint are short circuits to ground, and opens. The former tend to result in excessive currents, the latter, excessive voltages. Very fast acting protection in the terminal power bays is required to cope with these, and series electron tube regulators provide the means.

Terminal Plan

The plan of design of transmission terminal equipment, shown in Fig. 2, was based on the following objectives:

- (a) Use of standard, or modified standard equipment as far as possible.
- (b) To facilitate supply and maintenance of standard equipment, use of Bell System equipment in North America and Post Office equipment in the United Kingdom.
- (c) Provision of full duplicate equipment and means for quick shift between regular and alternate.
- (d) Provision of special equipment to fit Canadian requirements.
- (e) Provision for ample order-wire (speaker and printer) equipment, partly because there is no alternate undersea route.
- (f) Provision of no automatic loss regulation, because the loss changes are so slow; provision of three link pilot frequencies as a basis for manual regulation.

Much of the equipment is standard. This includes group modems, group connectors, pilot supply, frequency-shift telegraph equipment for printer circuits, etc. Supergroup modulators are modified standard coaxial carrier equipment. The channel modems for speaker and printer circuits are a combination of Type-C open-wire carrier active equipment,

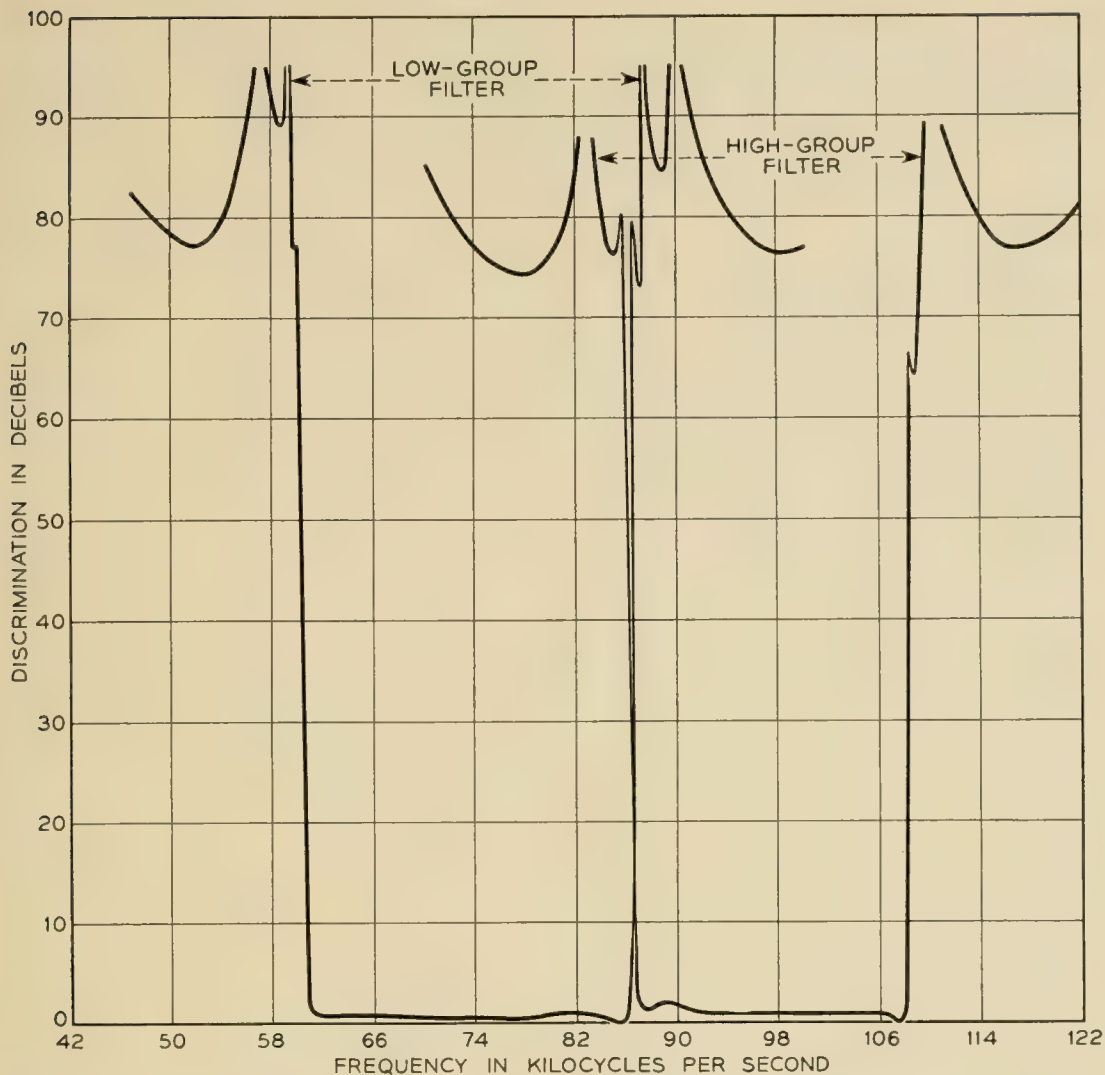


Fig. 11 — Splitting filters — loss frequency characteristic.

and filters designed for a military carrier system. The transmitting output amplifier is a modified version of one used in the Havana-Key West submarine cable system terminals; the modifications include broader frequency band, lower modulation and sharper overload cutoff. As in the coaxial carrier system, hybrid coils are provided at points in the Bell equipment where quick patching is needed.

The agreed Canadian bandwidth quota is 26 kc,¹ corresponding to six and one-half message telephone channels. The split between the $6\frac{1}{2}$ Canadian channels and the remaining $5\frac{1}{2}$ U.S. channels in the "split group" was accomplished by specially-designed group-frequency splitting filters, with very sharp cutoff. The loss-versus-frequency characteristics of these filters are shown in Fig. 11. The resulting Canadian half telephone channel next to the cut-apart was broad enough to accommodate

twelve frequency-shift telegraph channels with 120-cycle carrier telegraph spacing. It is planned to use one of these for automatic regulation of the other eleven. The U.S. half channel is at present unassigned.

Two telephone order wires (speakers) and two teletype order wires (printers) are provided for maintenance and administration purposes.

The two speakers occupy a 4-kc telephone channel which is the lower sideband of 20 kc on the line. The 4-kc band is divided into halves by Bell Type-EB split-channel equipment, which yields two narrow-band telephone channels. One of these is reserved for a "local" Clarendville-Oban circuit, available at all times to the personnel at these terminals. It constitutes a cleared channel which could be of great benefit in emergencies which might arise. The other narrow-band telephone channel, the "omnibus speaker", is extended through to other control points including the system metropolitan terminals.

The two printer circuits are frequency-shift voice-frequency telegraph channels, occupying line frequencies just below 16 kc. These are both brought to dc at Clarendville and Oban, and extended by carrier telegraph to the metropolitan terminals. One is normally a "through" circuit with teletypewriters at the metropolitan terminals only, and the other is an "omnibus" circuit with teletypewriters connected at Oban, Clarendville and other control points, permitting message interchange between all connected points.

92-kc group frequency pilots are transmitted over the Clarendville-Oban link only, and blocked at its ends. The line frequencies corresponding to the 92-kc group frequencies are 52, 100 and 148 kc. These pilots are used for manual regulation and maintenance of the Clarendville-Oban link. Information derived from them governs the manual setting of the temperature equalizers. The sending power of the 92-kc pilots is regulated to within ± 0.1 db variation with time.

84.080-kc group frequency pilots also appear on this link. These, however, are system pilots applied at the metropolitan terminals. They are useful for evaluation of overall system performance and for quickly locating transmission troubles.

Alarms are provided in the manner usual for multi-channel telephone systems. A special feature of the pilot alarm system is the "alarm relay channel", a third voice-frequency carrier telegraph channel just below the printer channels in line frequency. If all three eastbound 92-kc pilots fail to reach Oban, a signal is transmitted over the westbound cable which brings in a major alarm at Clarendville, and vice versa. This major alarm would provide news of a cable failure immediately to the transmitting cable terminal station. Failure of the pilot supply would, of course, give

the same alarm. The alarm relay channel circuit is arranged so that if it fails it does not give the major alarm.

A special alarm system for the cable current supply is described in the companion paper on that equipment.¹⁰

Fault Location

Length and inaccessibility have always imposed difficult requirements on fault locating techniques for undersea cable systems. Before the advent of repeatered systems, much effort was expended on use of impedance methods — i.e., those involving resistance and electrostatic capacitance measurements to locate a fault to ground, or a break.* Some work has been done also on magnetic pickup devices towed along the route by a ship.

With the advent of undersea repeaters, the problem has become even more difficult, both because of the effect of the repeater elements on impedance measurements, and because a fault may interrupt transmission without affecting the dc loop.

It has been necessary, therefore, to reassess the whole problem of fault location.

Location of Faults Affecting the DC Power Loop

Three types of measurement are made on submarine cables at present:

(a) Center conductor resistance, assuming a short circuit or a sea water exposure at some point, which might be either a fault or a break.

(b) Dielectric resistance, assuming a pinhole leak through the dielectric and a sea water exposure.

(c) Electrostatic capacitance, assuming a center conductor break insulated from sea water.

By these methods, faults and breaks in non-repeatered cables can be located although the accuracy of determination, and indeed the practical success of the operation, is dependent on the nature of the fault, the situation with respect to earth potential, the skill of the craftsman and many other factors.

The presence of cold repeaters in the circuit adds considerably to the difficulty, both because of the resistance/temperature characteristic of the electron tube heaters (153 in series with each cable) and because of polarization effects exhibited by the castor oil capacitors in the trans-

* In the literature, "fault location" is a generalized term encompassing the field. A "fault" is an exposure of the cable conductor to the sea without a break in the conductor. A "break" is an interruption of the conductor with or without exposure to the sea.

atlantic repeaters. These capacitors have a storage characteristic which varies with the magnitude of applied voltage, duration of its application and the temperature.

Because of this, the usual methods of fault location are expected to give results of doubtful accuracy and so a new approach is being made to the problem. The results of the work to date are beyond the scope of this paper.

Location of Transmission Faults

The previous section deals with the situation where the dc power loop has been opened or disturbed. Of equal importance is the location of transmission faults when the power loop is intact.

For this purpose, use is made of the discrete frequency crystal provided in each repeater.³ This crystal is effectively in parallel with the feedback circuit of the amplifier, and at its resonant frequency the amplifier has a noise peak of the order of 25 db, with an effective noise band of about 4 cycles. When one repeater fails, it is possible to recognize the noise peak of each amplifier from the receiving end back to the failed unit. The noise peaks from the faulty unit and all preceding amplifiers will be missing, indicating location of the trouble. Noisy amplifiers can be singled out by an extension of the technique.

Testing and Maintenance

Routine testing and maintenance activities at Oban and Clarenville can be divided into four parts: for the carrier terminals, for the station power equipment, for the cable current power units and for the high frequency line. Usual methods apply and the usual types of existing test equipment were provided with one exception.

The exception is a newly designed transmission measuring system. The system employs a decade type sending oscillator with continuous tuning over the final 1-kc range. Provision is built-in for precise calibration of output level. The receiving console contains a selective detector functioning as a terminating meter of 75-ohm input impedance, and can measure over the range - 120 to 0 dbm. Coils are supplied to permit measuring over 135-ohm circuits. The detector can be calibrated for direct reading when used as a terminating meter.

The system covers a frequency range of 10 kc to 1.1 mc. It is useful, therefore, for measuring much of the standard British and U.S. designed carrier terminal equipment as well as for normal line transmission measurements over both the North Atlantic and Cabot Strait (Newfoundland-Nova Scotia) submarine lines.

A special transmission measuring set is provided in the receiving console which facilitates measurements in the crystal frequency region 166–175 kc. This set is used for evaluating the performance of individual repeaters from shore. As an oscillator, it is capable of delivering an output of up to +8 dbm into a 75-ohm load, with exceptional frequency stability and finely-adjustable, motor-driven tuning. This is useful in locating and measuring the narrow-band response peaks of individual repeater crystals. The oscillator frequency is varied in the region of a particular peak, and the received power is measured at the peak frequency and at nearby frequencies. At the peak frequency, the crystal removes nearly all feedback in the repeater. Changes in repeater internal gain can thus be determined from shore.

As a detector, the set can measure from -110 to -60 dbm in a bandwidth of about 2 cycles. This enables it to measure crystal noise peaks which may be spaced as closely as 50 cycles apart. To reduce the random variations in such narrow-band noise, a “slow integrate” circuit, of the order of 10 seconds, is provided. Thus the crystal noise peaks can be compared in magnitude with the system noise level at closely adjacent frequencies.

ASSEMBLY AND TEST OF SYSTEM

General

Assembly and initial testing of the North Atlantic link occupied a span of about three years. Because of the geographical and political factors involved the job was a difficult one and required close cooperation among individuals in many different organizations on both sides of the ocean.

Clarenville

The cable station at Clarenville houses the carrier terminal, cable terminating and power equipment at the junction of the North Atlantic and Cabot Strait links. It was designed by United States architects from requirements furnished by A. T. & T. engineers and was approved by the other parties to the enterprise. The building was constructed by a Canadian firm under the supervision of a Canadian architect. The equipment was put in place and connected by installers of the Northern Electric Company and tested by representatives of Bell Telephone Laboratories, Eastern Telephone and Telegraph Company, Northern Electric Company and the British Post Office.

As Newfoundland is an island, shipments of equipment and supplies

to Clarenville were carried by sea or air to St. John's where they were trans-shipped by rail to Clarenville.

Oban

The cable station at Oban was executed by British contractors from designs drawn up by the British Post Office. Its equipment was installed by a British electrical contracting firm. Two Western Electric Company installers were present at Oban during the installation of the American-made cable terminating equipment and cable current supply bays to provide necessary liaison and interpretation of drawing requirements. The Oban equipment was tested by representatives of the British Post Office, Bell Telephone Laboratories and the firm that did the installing.

Here too, it was necessary to ship by boat or air, with subsequent trans-shipment by train and by truck.

Undersea Link

Perhaps the most interesting phase of all was the handling of the undersea section. The actual laying is described elsewhere, but much effort was required before the cable ship ever left the dolphins at the cable manufacturing plants.

Most of the cable was manufactured at the plant of Submarine Cables, Ltd., at Erith, on the Thames about fifteen miles downstream from London. A smaller quantity of cable was manufactured by Simplex Wire and Cable Corporation at its plant in Newington, N.H. As the cable was completed it was coiled in repeater section lengths in huge tanks. The ends of each repeater section were left available at a "splicing platform" where the repeaters, manufactured at Hillside, N.J., were spliced in. Spliced repeaters were located in water filled troughs for protection against overheating while testing.

All repeaters were armored by Simplex at Newington. This involved their transportation by truck from the Western Electric shop at Hillside where the flexible repeaters were manufactured, to Newington. While this sounds simple it was actually a very carefully planned and controlled operation because of the need to avoid subjecting the units to any but the most necessary and unavoidable hazards. Consequently, the truck used for the purpose was specifically selected for size and construction and was provided with a heating unit in the body. The route was carefully surveyed in advance for unusual hazards and the truck speed was limited to a very modest value.

After the repeaters were armored, those to be incorporated in Simplex-made cable were spliced at Newington. Those destined for application in

British-made cable were shipped by truck, following the same precautions, to Idlewild Airport in New York. From there they proceeded by special freight plane — one or two at a time — to London Airport. Here they were again trans-shipped by truck to Erith.

The repeaters were housed in shipping cases¹² of a very strange shape and considerable size. These cases were provided with max-min thermometers and with impactograph devices that would record the maximum acceleration to which the repeaters had been subjected.

When ship loading time arrived, the cable and its contained repeaters were transported over a system of sheaves from the tanks of the cable manufacturer to the tanks on shipboard. This process offered no particular obstacle so far as the cable itself was concerned, but the requirement against unnecessary bending of the repeater meant that a great deal of special attention had to be given to avoiding unnecessary deviations from a straight run, and to lifting the repeaters around sheaves where the direction of the loading line changed materially. Auxiliary protection was provided on each repeater in the form of angle irons with flexible extensions on each end which served to restrict any bending to the core tube region of the repeater and safely limited the magnitude of the bends in this region.³ One further precaution was observed — that of energizing the repeaters during the loading process. This accomplished two things. First, it permitted continuous testing; second, it reduced the hazards of possible damage to the electron tubes during loading, as the tungsten heaters are much more ductile when hot and the glassware of the tube is more resistant to shock and vibration.

SYSTEM PERFORMANCE

General

At the time of this writing, the No. 2 cable has just been laid, but the No. 1 cable has had eleven months of successful life under test, a pre-service period of probably greater duration than has been granted to any land system of comparable length and cost. It has been under constant observation and test, largely by the people who will operate it when it goes into service. The pre-service measurements have been much more extensive, in quantity and scope, than the routine tests which will be made after the system goes into service.

Net Loss Tests

Net Loss versus Frequency

Results of measurements in June, 1956, on the No. 1 cable are shown in Figs. 12 and 13. The former covers the equalized high-frequency line.

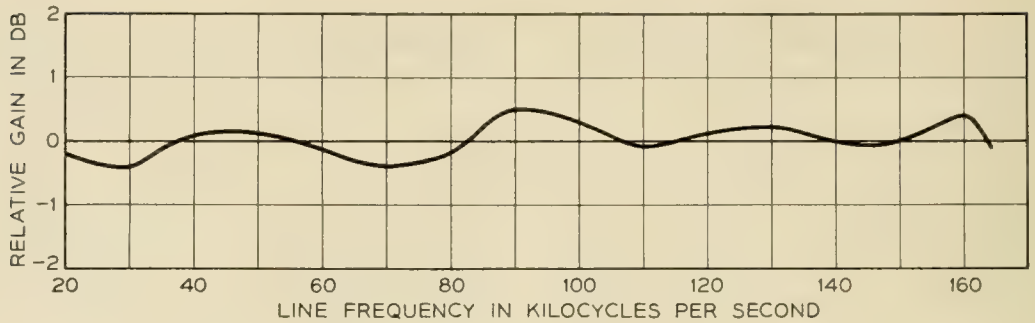


Fig. 12 — Gain versus line frequency — No. 1 cable.

Fig. 13 gives the corresponding group-to-group frequency response of the deep-sea link. These frequency characteristics will vary a little from time to time, depending partly on the change since the last adjustment of the temperature equalizer.

Note that while the slope in gross cable loss between 20 and 164 kc is about 2,100 db, the net loss of the equalized high frequency line over this band varies only about 1 db.

Net Loss versus Time

The net loss of the undersea system (cable plus repeaters) has decreased slowly since the system was laid. The decrease is approximately cable shape, proportional to the square root of frequency. In a year it has amounted to about 5 db at the top frequency. In the early months the change was almost directly proportional to time, but later the rate slowed, as shown in Fig. 14.

The changes shown in the figure are due partly to change in cable temperature and partly to slow aging of the cable. Detailed studies of repeater crystal frequency changes have resulted in only an approximate separation into temperature effects and aging. However, it seems reasonable to assume that at the end of a one-year cycle there would be little

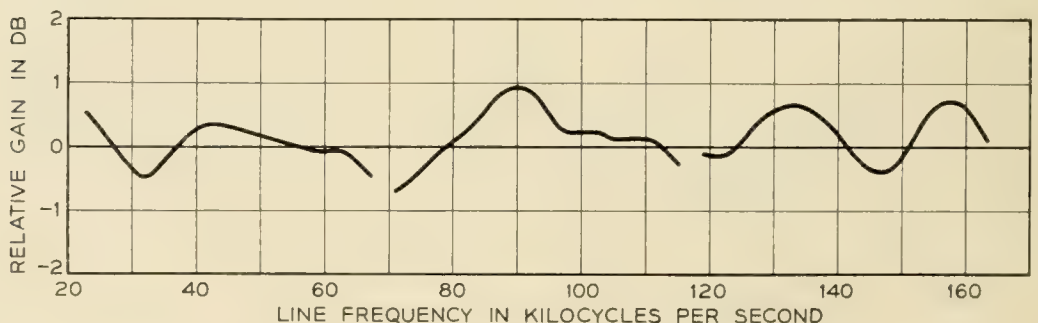


Fig. 13 — Group to group frequency response — No. 1 cable.

if any net change in the cable temperature averaged over the whole cable length. The greatest possibility of change would be in the shallow-water sections, and crystal measurements indicate that at the end of a year the average temperature of the shallow-water sections returned nearly to its initial value. Hence about 5 db seems chargeable to one year's aging. Extrapolation into the future, however, is uncertain. Theoretical considerations lead to the idea that the rate of cable aging should decrease.

Information on long-term change in transmission loss of previous cables is very limited. The Havana-Key West cables were accurately measured just after laying, and again five years later. The change in that system due to aging is very small, if indeed there is a change. That cable is generally like the transatlantic, though it is smaller and has a perhaps significant difference in construction of the central conductor.

The above applies to the net loss of the undersea part of the system only. The group-to-group net loss variation with time of the Clarendville-Oban link has been held within a much smaller range by temperature equalizer adjustments. Practices have been worked out which it is believed will hold the in-service variation with time to a fraction of a db.

Net Loss versus Carrier Frequency Power Level

Single-frequency tests of carrier frequency output power versus input power were made, up to a test power a little below the estimated overload of the highest-level repeater. An increase of 0.1 db in system net loss occurs at a power 2 to 3 db below the load limit of the transmitting amplifier. This is about 15 or 16 db above the expected rms value of the in-service system busy hour load. The 0.1 db change is presumably due to the cumulative effect of smaller changes in several of the undersea amplifiers.

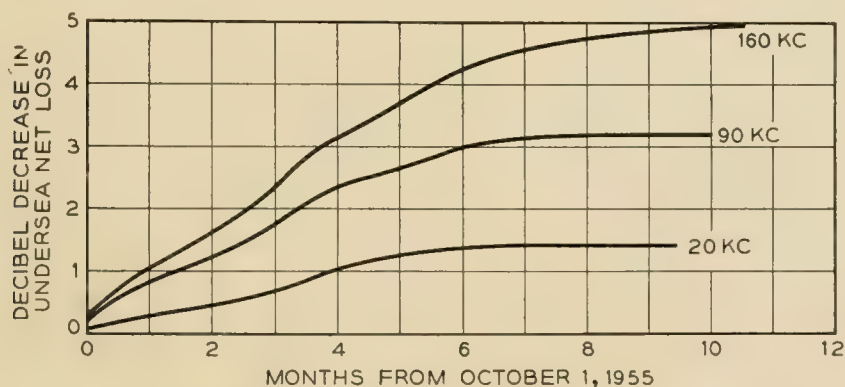


Fig. 14 — Change in system gain in first eleven months.

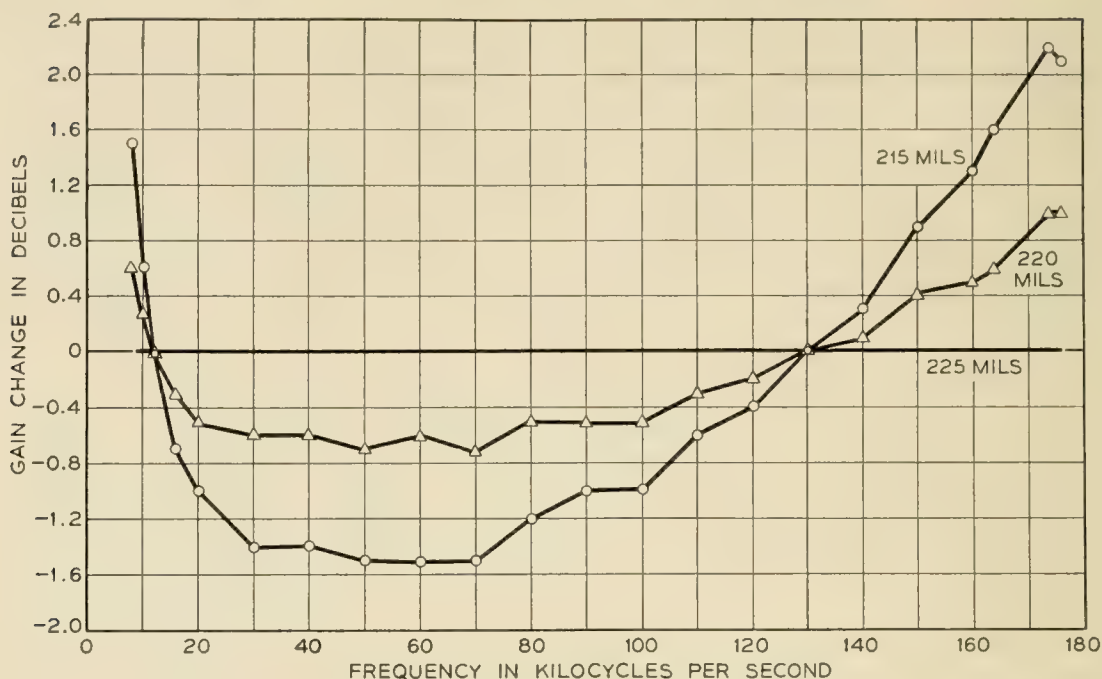


Fig. 15 — Effect of cable current on system gain — No. 1 cable.

Net Loss versus DC Cable Current

Changes in cable current affect the repeater gain to a slight extent, the amount depending on the magnitude and phase of the feedback in the repeater as a function of frequency. Measured changes in the loss of the 2,000-mile system, for currents of 5 and 10 milliamperes less than the normal value of 225 milliamperes, are shown in Fig. 15. Under normal conditions the automatic control will hold the cable current variation within ± 0.5 milliampere.

The shape of the curve on Fig. 15 is almost the same as that computed in advance from laboratory measurements on model repeaters.

System Noise

Shown in Fig. 16 are values of noise on the No. 1 cable system measured in the Fall of 1955 and again in the Spring of 1956. The noise increase is compatible with the decrease in undersea system net loss during this period. To prevent overload, the loss in the transmitting temperature equalizer has to be increased as the undersea loss decreases; this lowers the levels of the various parts of the undersea system by various amounts.

The noise shown in the top channel exceeds the 36 dba objective by a small amount. The excess can be recovered, if necessary, by certain changes in the terminals without recourse to compandors.

Fig. 16 shows also the noise on the No. 2 cable system shortly after

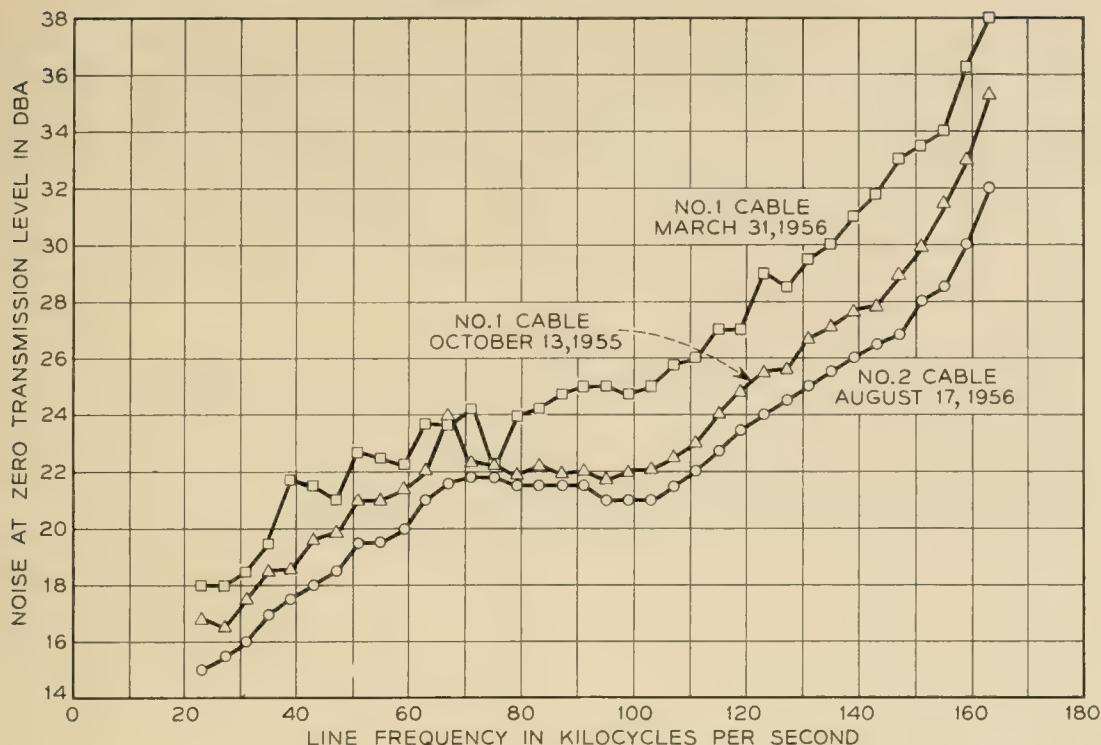


Fig. 16 — System noise.

completion of laying. The noise is lower than on the No. 1 cable. This is because results of experience on the No. 1 cable were utilized in better choice of cable repeater section lengths and in better equalization while laying the 1956 cable.

Modulation Tests

Two-tone modulation tests of second and third order products were made, using a large number of successive frequency combinations. The highest level modulation products were at least 60 db below the 1-milli-watt test tones, at zero transmission level. This is approximately the value computed before the system was laid. Most of the modulation products were down substantially more than 60 db. Probably various causes contributed to the good performance, including the effect of misalignment in lowering repeater levels, and small propagation-time differences which minimize in-phase addition of third-order products from successive repeaters.

Telegraph Transmission Tests

At present writing, telegraph tests have been made only on the printer (telegraph order wire) channels, and without a system multi-channel

load. With the proposed specific telegraph level (STL) of -30 db (i.e. telegraph signal of -30 dbm per telegraph channel at 0 db telephone transmission level), the telegraph distortion was too low to measure reliably; it was possible to send clear messages under test conditions with a signal 36 db weaker than this.

Crystal Tests

All of the peaks of noise at the crystal resonance frequencies were easily discernible.

The crystal gain values all lay in the range from 23.6 to 27.2 db, with 60 per cent of them lying in the range from 25 to 26 db. Crystal gain, as used here, is the difference between the system gain at a repeater crystal frequency and the average of the gains at 50 cycles above and below this frequency; the latter value is approximately the gain if the crystal were absent. No significant changes in crystal gain have occurred in eleven months of system life, and none were expected. These measurements are to be continued over the years, as an indication of electron tube aging, as explained in the companion paper on the repeater.³

A series of measurements of the frequency of each of the 51 repeater crystals versus time has been made. (Any of these frequency determinations can be checked on the same day within ± 0.1 cycle with the special test apparatus and techniques used.) The crystals were designed to be extremely stable in frequency, so that measurements on one repeater, made at the land terminal, would not be affected by the combined effect of 50 other crystals at 100 -cycle spacings. The crystal frequency varies only about $-\frac{1}{2}$ cycle per degree Fahrenheit increase in sea-bottom

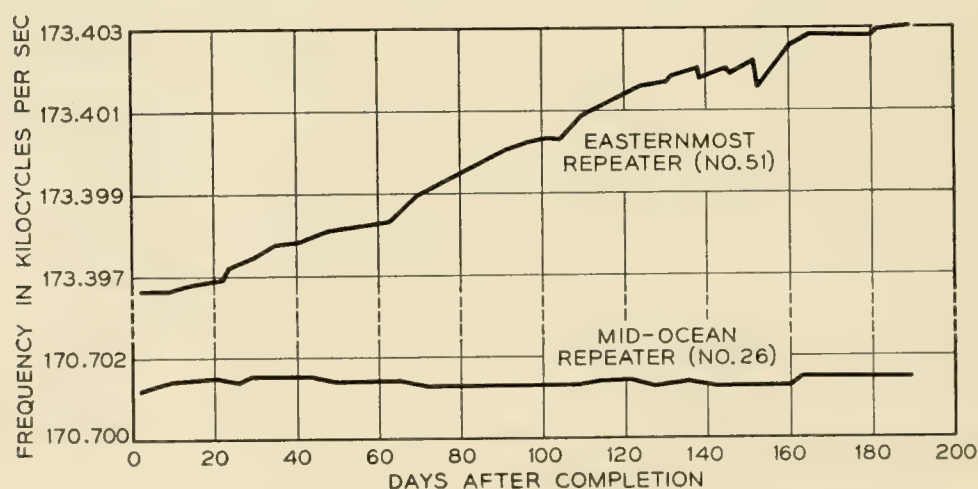


Fig. 17 — Resonant frequency versus time — shore and deep sea crystals.

temperature. The change in frequency due to crystal aging in 11 months under the sea is considered to lie in the range 0 to -0.4 cycle.

Although the crystals were not designed as sea-bottom temperature indicators and are within a repeater housing, it seems likely that with precise techniques and, after some further stabilization, some uniquely accurate information on change of sea-bottom temperature will be obtainable.

In accordance with previous oceanographic knowledge of sea-bottom temperature, the frequency changes have been larger in the crystals near shore. The greatest change has been in the repeater nearest Oban. This is about 17 miles from shore and in water only about 50 fathoms deep. Its frequency change, and that of a typical deep-sea repeater, are shown in Fig. 17.

FUTURE CONSIDERATIONS

Spare Equipment for Cable Stations

In a system as far-flung geographically as this one, and as important, much thought must be given not only to the supplies needed for routine replacement of expendable items, but also to major replacements necessitated by fire or other causes. Accordingly, a schedule of spare equipment has been established, divided into two groups — “shelf” and “casualty”.

Shelf spares are carried in the station itself, and include items such as dial lights, electron tubes and dry cell batteries, which have limited life in normal service. These are maintained in the cable station in quantity estimated as adequate for two years of operation.

Casualty spares embrace those major, essential frames and equipments without which the system cannot be operated. There are two subdivisions: those which are in common use in the telephone plant and are, therefore, available by “cannibalization”, and those which are special to the project, like the cable current supplies and the cable terminating bays. The common-use items are not stocked as casualty spares. Spares of the special items have been built and are stored in locations remote from the cable stations. In event of catastrophe, these can be drawn out and flown as near the point of need as possible, for onward-shipment by available means.

Spare Equipment for the Undersea Link

Although every effort has been made to produce a trouble-free system, the underwater link is still subject to the hazards of trawlers, icebergs,

anchors and submarine earthquakes or land slides. There must be considered, too, the possibility of fault from human failure.

So it is necessary to contemplate the replacement of a length of cable, or a repeater or equalizer. For such contingency, spare cable of the various armor types has been stored on both sides of the Atlantic. Spare repeaters and equalizers are also available.

In addition, a spare called a "repair repeater" has been stocked. Need for this arises from the fact that except in very shallow water, a repair cannot be effected without the addition of cable over and above the length which was in the circuit initially. The amount of cable which must be added is a function of water depth, condition of the sea at the time of the repair, and the amount of cable slack available in the immediate vicinity of the point in question.

When excess cable is introduced in amount sufficient to significantly reduce the system operating margins, its loss must be compensated. Hence, the repair repeater.

A repair repeater is a two-tube device, essentially like a regular repeater although its impedances are designed to match the cable impedances at input and output ends. However, its gain is sufficient to offset only about 5.3 miles of cable. A second type of repair repeater is under consideration, to compensate for about 15 miles of cable.

Long-Term Aging

General

In a system with some 3,200-db gross loss at its top frequency between points which are accessible for adjustment, a long-term change in loss of only one per cent would have a profound effect on system performance.

For this reason, the repeater design included careful consideration of net gain change over the years.³ The degree of control over aging is such that in a period of at least 20 years, and perhaps much longer, the estimated change in 51 repeaters might total 8 db added gain at the top frequency. The gain variation with frequency would be proportional to either curve of Fig. 15, and the rate of aging would be slower in earlier than in later years.

Estimates of cable aging are discussed in the section on "Net Loss Tests."

Means of Combatting Aging

The effects of aging would become important on the top channels first. Remedial measures to improve signal-to-noise, especially in the top

channels, include: possible increase of transmission level at input of the final transmitting and load limiting amplifier in the transmitting terminal; pre-distortion ahead of this amplifier; compandors; increase of dc current; and undersea re-equalization in later years. This last would be very expensive, and so it is necessary to examine fully the possibilities of the other measures.

The penalty for increasing the transmission level at the input of the transmitting amplifier is more peak-chopping and modulation-noise peaks. The improvement that could be realized in this way is probably fairly small.

Pre-distortion is accomplished by inserting ahead of the transmitting amplifier a suitable shaping network adjusted for gain in the top part of the band and loss at lower frequencies. A complementary network (restoring network) is placed at the receiving terminal. This measure would improve signal-to-noise in the uppermost part of the band and reduce it in lower channels which have less noise. Some 3-db improvement might be thus realized in the top channel.

Compandors would give an effective signal-to-noise improvement of up to about 15 db for message telephone service, but none for services such as voice-frequency telegraph. Compandors would be applied only to those channels needing them. They halve the range of talker volume, but also double the transmission variations between compressor and expander, and thus tend to require some increase in the overall channel net loss. The program (music) channels are already equipped with compandors which use up a part of the obtainable advantage.

If the combination of such measures netted an effective message signal-to-noise improvement of 20 db in the top channel, this would counter-balance aging of some 28 db in this channel, if the aging were uniformly distributed along the system length. Thus considerable aging could be handled without undersea modification.

ACKNOWLEDGEMENTS

A system of the complexity of the one described obviously results from teamwork by a very large number of individuals. However, no paper on this subject could be written without acknowledgement to Dr. O. E. Buckley and J. J. Gilbert and O. B. Jacobs, now retired from Bell Telephone Laboratories. All of the early and fundamental Bell System work on repeatered submarine cable systems, and the concept of the flexible repeater, came from these sources and from their co-workers. Messrs. Gilbert and Jacobs have also contributed to the present project.

REFERENCES

1. E. T. Mottram, R. J. Halsey, J. W. Emling and R. G. Griffiths, Transatlantic Telephone Cable System — Planning and Over-All Performance. See page 7 of this issue.
2. A. W. Lebert, H. B. Fischer and M. C. Biskeborn, Cable Design and Manufacture for the Transatlantic Submarine Cable System. See page 189 of this issue.
3. T. F. Gleichmann, A. H. Lince, M. C. Wooley and F. J. Braga, Repeater Design for the North Atlantic Link, See page 69 of this issue.
4. J. J. Gilbert, A Submarine Telephone Cable with Submerged Repeaters. B.S.T.J., **30**, p. 65, Jan., 1951.
5. C. W. Carter, Jr., A. C. Dickieson and D. Mitchell, Application of Companders to Telephone Circuits, A.I.E.E. Trans., **65**, pp. 1079–1086, Dec. Supplement, 1946.
6. B. D. Holbrook and J. T. Dixon, Load Rating Theory for Multi-Channel Amplifiers. B.S.T.J., **18**, p. 624, July, 1939.
7. W. R. Bennett, Cross Modulation in Multi-Channel Amplifiers, B.S.T.J., **19**, pp. 587–610, Oct., 1940.
8. R. A. Brockbank and C. A. Wass, Non-Linear Modulation in Multi-Channel Amplifiers, J.I.E.E., March, 1945.
9. J. S. Jack, Capt. W. H. Leech and H. A. Lewis, Route Selection and Cable Laying for the Transatlantic Cable System. See page 293 of this issue.
10. G. W. Meszaros and H. H. Spencer, Power Feed Equipment for the North Atlantic Link. See page 139 of this issue.
11. J. O. McNally, G. H. Metson, E. A. Veazie and M. F. Holmes, Electron Tubes for the Transatlantic Cable System. See page 163 of this issue.
12. H. A. Lamb and W. W. Heffner, Repeater Production for the North Atlantic Link. See page 103 of this issue.
13. P. T. Haury and L. M. Ilgenfritz, Air Force Submarine Cable System, Bell Lab. Record, **34**, pp. 321–324, Sept., 1956.

Repeater Design for the North Atlantic Link

T. F. GLEICHMANN,* A. H. LINCE,* M. C. WOOLEY* and
F. J. BRAGA*

(Manuscript received October 8, 1956)

Some of the considerations governing the electrical and mechanical design of flexible repeaters and their component apparatus are discussed in this paper. The discussion includes description of the feedback amplifier and the sea-pressure resisting container that surrounds it. Examples are given of some of the extraordinary measures taken to ensure continuous performance in service.

INTRODUCTION

Repeaters for use in the transatlantic submarine telephone cable system had to be designed to resist the stresses of laying, and to withstand the great pressures of water encountered in the North Atlantic route. In anticipation of the need for such a long telephone system in deep water, development work was started over 20 years ago on the design of a flexible repeater that could be incorporated in the cable and be handled as cable by conventional cable ship techniques. Successful completion, in 1950, of the design and construction of the 24-channel Key West, Florida-Havana, Cuba system,³ led to the adoption of similar repeaters designed for 36 channels for the North Atlantic link discussed in companion papers.^{1, 2}

Repeater transmission characteristics determine, to a large extent, the degree to which system objectives can be met. In this repeater, significant characteristics are:

(a) *Noise and Modulation.* These were established by the circuit configuration and by the use of the conservative electron tube⁸ developed for the Key West-Havana project.

(b) *Initial Misalignment,* or mismatch of repeater gain and cable loss throughout the transmitted band of frequencies. A match within 0.05 db was the objective. This affected both the design and the precision required in manufacture.

* Bell Telephone Laboratories.

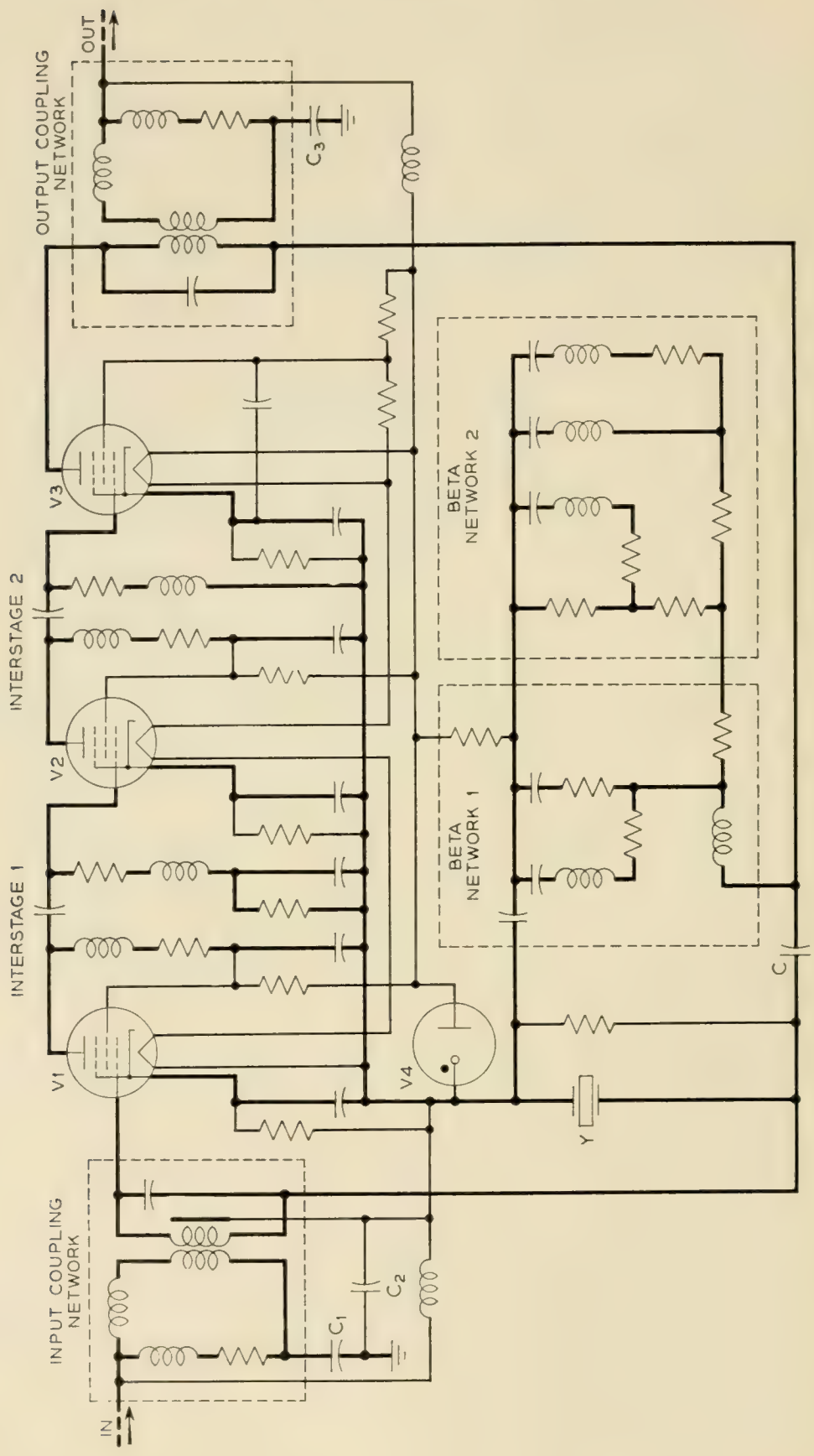


Fig. 1 — Repeater schematic.

(c) *Aging.* As electron tubes lose mutual conductance with age, repeater feedback decreases, repeater gain changes, and misalignment is affected. Decrease in feedback increases the gain at the higher frequencies so that the signal input must be reduced to prevent overloading, resulting in a signal-to-noise penalty. Gain increase is inversely proportional to the amount of feedback; in these repeaters, 33 to 34 db of feedback was the objective to keep this source of misalignment in bounds.

Because repeaters are inaccessible for maintenance, facilities are provided to enable the individual repeater performance to be checked from the shore end. This feature also permits a defective repeater to be identified in the event of transmission failure.

REPEATER UNIT

The repeater, for the sake of discussion, may be divided into two parts, (1) the repeater unit, which contains the electron tubes and other circuit components and (2) the water-proof container and seals which house the repeater unit.

Circuit

The circuit of the repeater unit is shown in Fig. 1. It is a three-stage feedback amplifier of conventional design with the cathodes at ac ground. The amplifier is connected to the cable through input and output coupling networks. Each coupling network consists of a transformer plus gain-shaping elements and a power separation inductor.

The coupling networks directly affect the insertion gain as do the two feedback networks. The design of these networks controls the insertion gain of the amplifier. The required gain (inverse of cable loss) is shown in Fig. 2. The 39 db shaping required between 20 and 164 kc is divided approximately equally among the input and output coupling networks and the feedback networks.

The interstage networks are of conventional design. The gain of the first interstage is approximately flat across the band. The second interstage has a sloping characteristic, the gain increasing with frequency. The gain shaping of these networks offsets the loss of the feedback networks so that the feedback is approximately flat across the band.

Plate and heater power is supplied to the repeater over the cable.⁴ The plate voltage (approximately 52 volts) is obtained from the drop across the heater string. The dc circuits are isolated from the container by the high voltage blocking capacitors C_1 , C_2 and C_3 .

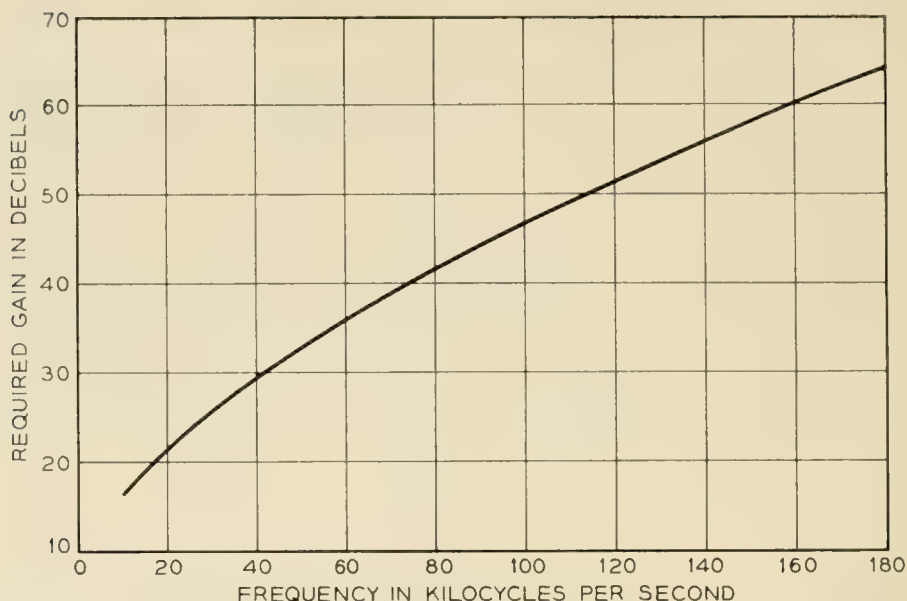


Fig. 2 — Required insertion gain.

Gain Formula

The circuit of Fig. 1 may be represented by a simplified circuit consisting of an input coupling network, a three stage amplifier, an output network and a feedback impedance Z_β as shown in Fig. 3. From this figure it can be shown that the insertion gain of the repeater is given by:⁶

$$e^\theta = \frac{2e^{\theta_1}e^{\theta_2}Z_c}{Z_\beta} \left[\frac{\rho_i \rho_0 g_{mT} Z_1 Z_2 Z_\beta}{1 - \rho_i \rho_0 g_{mT} Z_1 Z_2 Z_\beta} \right] \quad (1)$$

when $Z_\beta \ll g_{mT} Z_p Z_g Z_1 Z_2 \gg (Z_0 + Z_p)$ and where

$$\rho_i = \frac{Z_0}{Z_i + Z_g} \quad \text{and} \quad \rho_0 = \frac{Z_p}{Z_0 + Z_p}$$

are “potentiometer terms”. The gain of the input network is defined as $e^{\theta_1} = V/E_i$; where V is the open circuit voltage of the input network with E_i as the source, and the gain of the output network is defined as $e^{\theta_2} = i/I_1$. This expression may be put in familiar form by recognizing that $\rho_i \rho_0 g_{mT} Z_1 Z_2 Z_\beta$ is $\mu\beta$, the feedback around the loop. Hence

$$e^\theta = \frac{2e^{\theta_1}e^{\theta_2}Z_c}{Z_\beta} \left[\frac{\mu\beta}{1 - \mu\beta} \right] \quad (2)$$

Equation (2) shows that the insertion gain of the repeater is the product of five factors, namely:

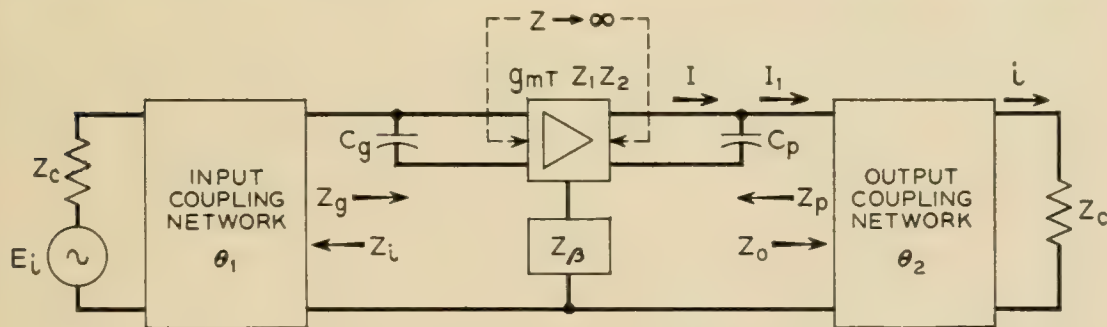
- (1) e^{θ_1} — the gain of the Input Network
- (2) e^{θ_2} — the gain of the Output Network
- (3) Z_c — the cable impedance
- (4) Z_β — the feedback impedance
- (5) $(\mu\beta/1 - \mu\beta)$ — the $\mu\beta$ effect term

It should be noted that a number of simplifying assumptions have been made. For example, the effect of grid plate capacitance has been neglected. In addition the β circuit has been assumed to be a two terminal impedance whereas it is actually a four terminal network. However, in the pass band and over a large part of the outband of the repeater these simplifications give a very good approximation to the true gain of the repeater.

In the pass band $(\mu\beta/1 - \mu\beta)$ is very nearly unity so that the gain controlling factors are e^{θ_1} , e^{θ_2} , and $1/Z_\beta$ assuming that Z_c is fixed.

Coupling Networks

The input and output networks are essentially identical. The networks are of unterminated design and therefore do not present a good termination to the cable at all frequencies which results in some ripple in the system transmission characteristic at the lower edge of the band and makes the repeater insertion gain sensitive to variations in the cable impedance. However, this arrangement has the advantage of maximum



V = OPEN CIRCUIT VOLTAGE OF INPUT COUPLING NETWORK
WITH E_i AS THE SOURCE

θ_1 = GAIN OF INPUT COUPLING NETWORK DEFINED AS $e^{\theta_1} = V/E_i$

θ_2 = GAIN OF OUTPUT COUPLING NETWORK DEFINED AS $e^{\theta_2} = i/I_1$

$Z_1 Z_2$ = INTERSTAGE IMPEDANCES

g_{mT} = PRODUCT OF g_m OF THREE AMPLIFIER TUBES

Fig. 3 — Simplified amplifier circuit.

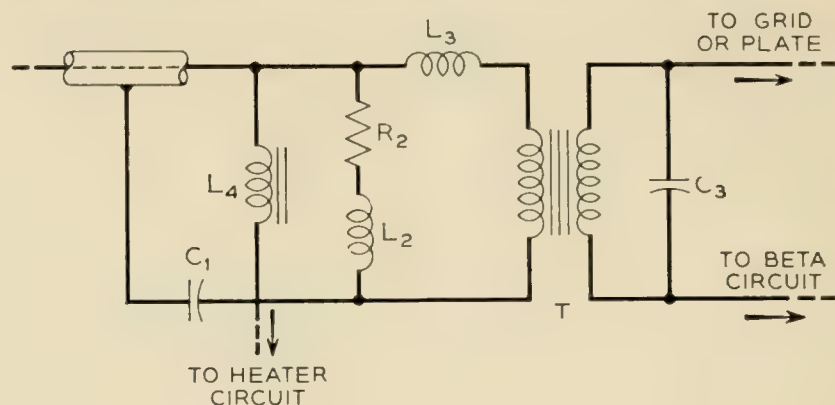


Fig. 4 — Coupling network.

signal to noise performance, highest gain, and most effective shaping with a minimum of elements. A minimum of elements is important in view of the space restrictions imposed by the flexible repeater structure. The sensitivity of the gain to variation in impedance is minimized by close manufacturing control of the cable and networks.

The schematic of a coupling network is shown in Fig. 4 and the equivalent circuit in Fig. 5. Capacitor C_1 and inductor L_4 are part of the power separation circuit. The effect of L_4 in the transmission band is negligible and it has been omitted from the equivalent circuit. However C_1 is in the direct transmission path and has a small effect at the lower edge of the band so that it becomes a design parameter. The combination R_2 , L_2 controls the low-frequency gain shaping of the network. Inductor L_3

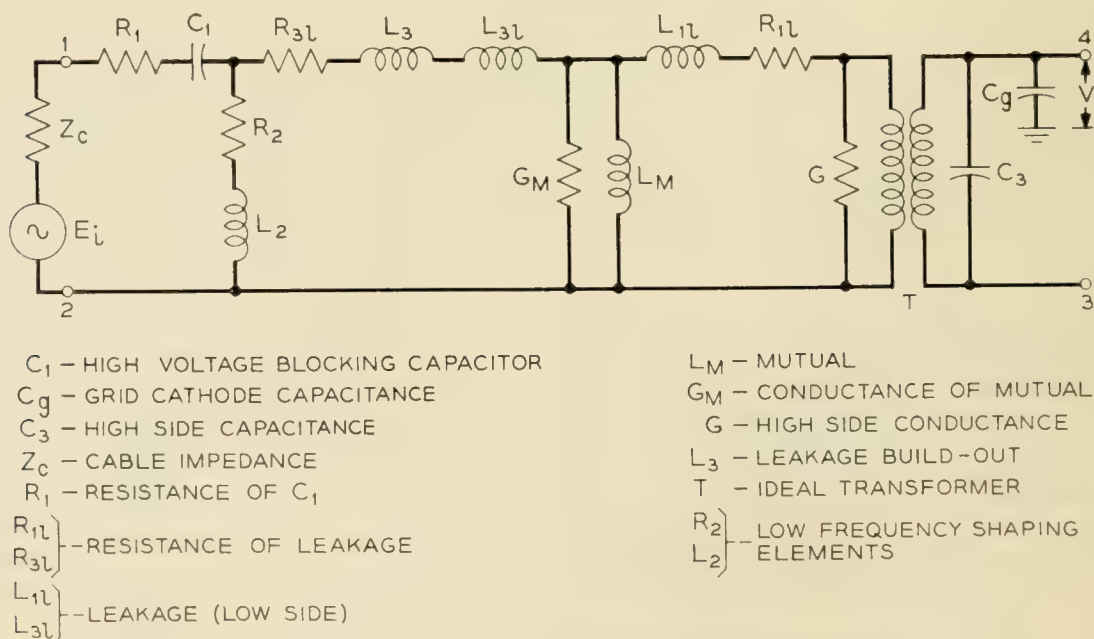


Fig. 5 — Equivalent circuit of coupling network.

builds out the leakage inductance of the transformer and together with capacitor C_3 controls the shaping at the top end of the band. These elements are adjusted during manufacture of the networks to provide the desired shaping.

The equivalent circuit is an approximation to the true transformer circuit. By standard network analysis techniques the ratio V/E_i , the gain of the network, can be obtained. The agreement between measurements and computation is sufficiently close, several hundredths of a db, to insure that the representation is good.

Each coupling network is designed to provide approximately one-third of the total shaping required, or 13 db. While these networks are

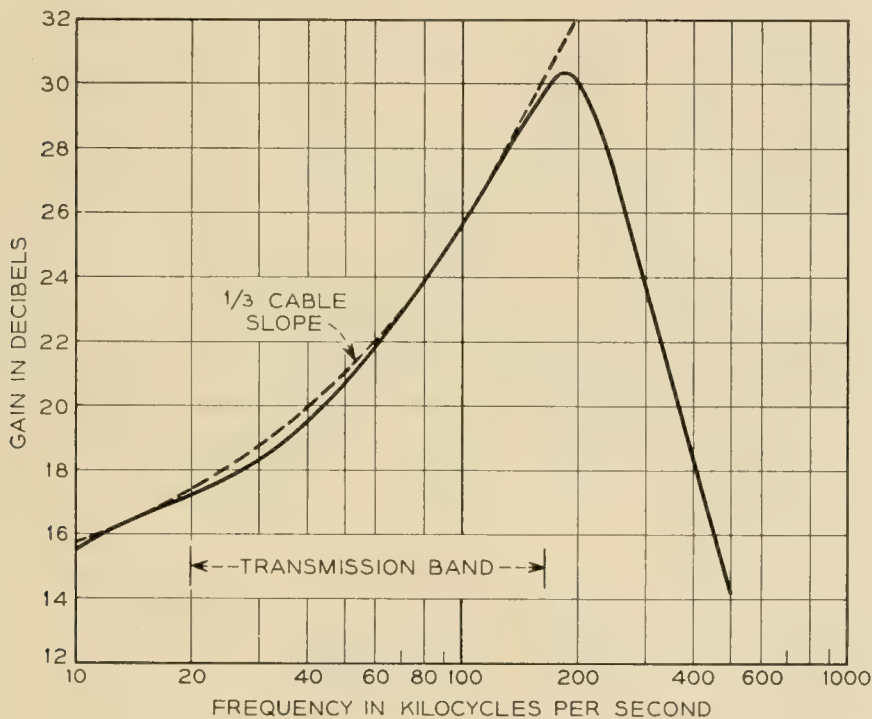


FIG. 6 — Gain of input coupling network.

outside the feedback path, the impedances which they present to the amplifier are important factors in the feedback design. It can be seen from Fig. 3 that at the amplifier input the proportion of the feedback voltage which will be effective in producing feedback around the loop is dependent upon the potentiometer division between the grid-cathode impedance of the first tube and the impedance looking back into the coupling network. The greater the gain shaping of the network, the greater the potentiometer loss. The maximum gain which can be obtained from the coupling network is limited by the capacitance across the circuit. This capacitance cannot be reduced without increasing the

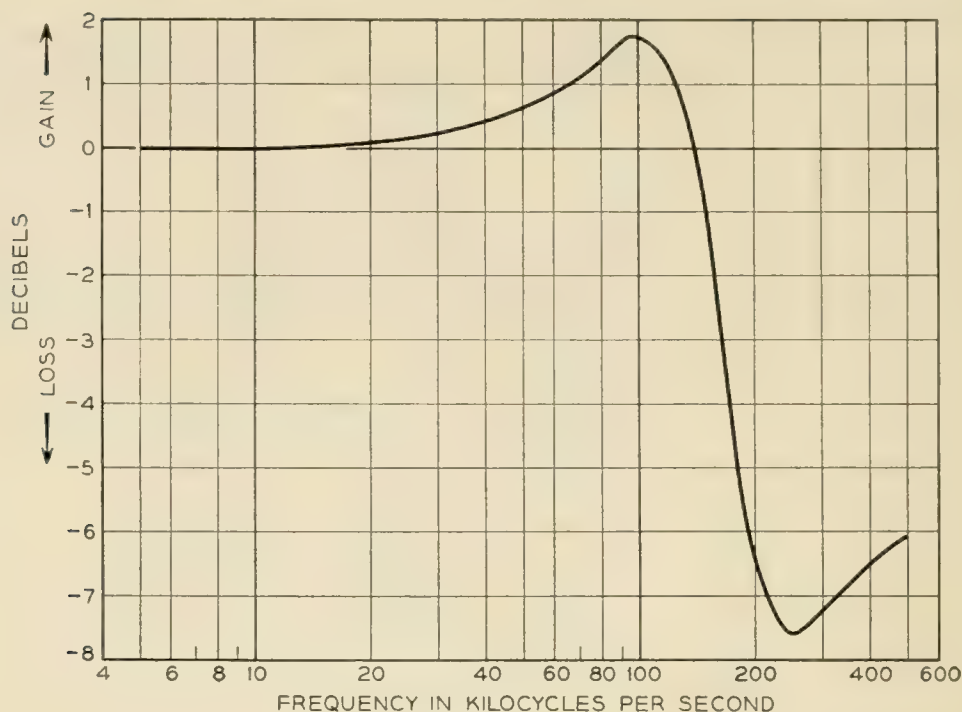


Fig. 7 — Input potentiometer term.

potentiometer loss, and seriously limiting feedback. In this design an acceptable compromise is made when the ratio of network capacitance to grid-cathode capacitance has been fixed at 1.2 as suggested by Bode.⁶ The gain through the input network and the deviation from one-third cable shape is shown in Fig. 6. A typical potentiometer term is shown in Fig. 7

Similar considerations apply to the output network with the further restriction that the impedance presented to the output tube should be about 40,000 ohms at the top edge of the band for optimum modulation performance.

The coupling networks have a temperature characteristic which must be taken into account in the insertion gain of the repeater. The characteristic is due to variations in the resistance of C_1 and R_2 with temperature. This amounts to 0.005 db per degree F at 20 kc, decreasing with frequency, becoming negligible above 80 kc.

Beta Circuit

The beta or feedback network is designed to complement the combined characteristics of the input and output coupling networks and mop-up residual effects, such as those due to $\mu\beta$ effect and coupling network temperature coefficients. The network also provides the dc path for the output tube plate current.

The configuration of the beta circuit is shown in Fig. 8. In the pass band it is a two terminal network whose impedance varies from about 300 ohms at 20 kc to 70 ohms at 164 kc. It consists of essentially two parts. The elements to the left of the dotted line provide the major portion of the shaping. With these the repeater is within ± 0.7 db of the required gain. The series resonant circuits to the right of the dotted line reduce this to the ± 0.05 db set as the objective.

The mopping up elements are connected to the main portion of the beta circuit through a resistance potentiometer R_1 , R_2 and R_3 . This scales the elements of the resonant circuits to values which would meet mounting space and component restrictions.

Built-in Testing Features

The crystal Y and capacitor C , Fig. 1, in the feedback path provide the means for checking the repeater from the shore station. The crystal is a sharply tuned series-resonant shunt on the feedback path which reduces the feedback at the resonant frequency and produces a narrow peak in the insertion gain characteristic of the repeater. The feedback reduction, and hence the peak gain, is controlled by the potentiometer divider formed by the reactance of the capacitor and the series resonant resistance of the crystal. The crystal and capacitor are chosen so that substantially all the feedback is removed from the repeater. With no feed-

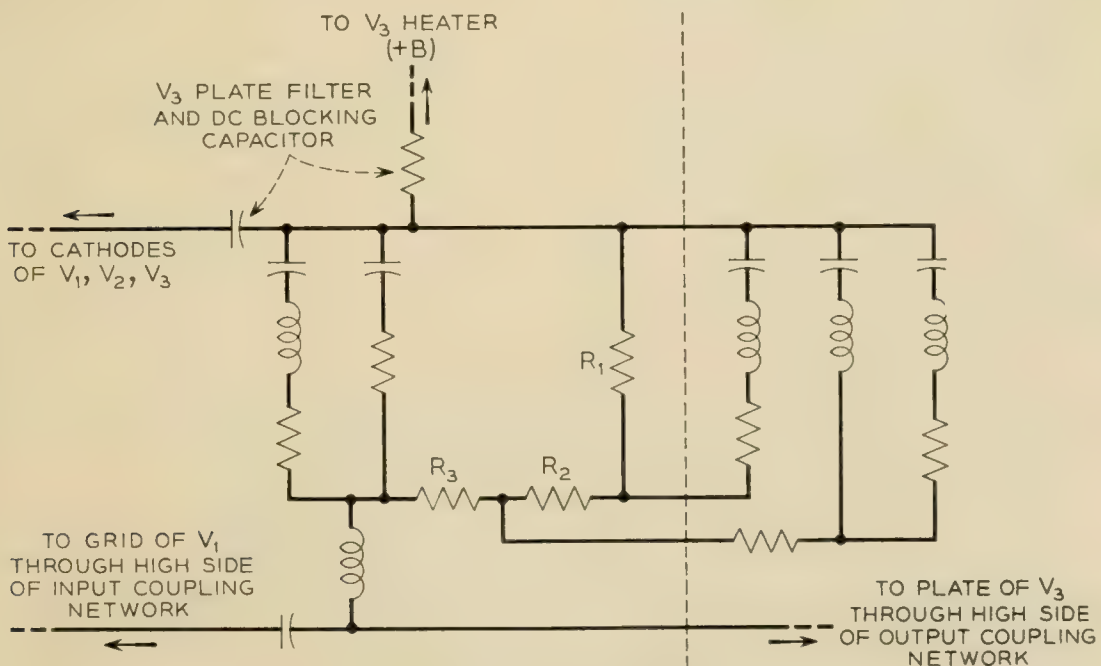


Fig. 8 — Beta network.

back the peak gain is proportional to the mutual conductance of the three tubes.

At frequencies well off resonance the impedance of the crystal is high so that no reduction in feedback results. Periodic measurements of gain at the resonant frequency relative to measurements made at a frequency off resonance will show any changes in the tubes. The crystal frequency is different for each repeater so that by measuring the gain from the shore stations at the various crystal frequencies it is possible to monitor the performance of the individual repeaters.

The increase in gain at the peak is approximately 25 db. The crystal frequencies, spaced at 100-cycle intervals, are placed above the normal transmitted band between 167 and 173.4 kc.

Thermal noise always present at the input to the repeater, is also amplified over the narrow band of frequencies corresponding to the peak gain in each repeater so that at the receiving end of the line there are a series of noise peaks, one for each repeater. Should a repeater fail, the noise peaks of all repeaters between the faulty repeater and the receiving end will be present and those from repeaters ahead will be missing. By determining which peaks are missing the location of the failed repeater can be determined. It is obvious that to locate a faulty repeater the power circuit must be intact. To guard against power interruption owing to an open electron-tube heater, a gas tube V4, Fig. 1, is connected across the heater string as a bypass.

Loop Feedback

The design of the feedback loop follows conventional practice. The restrictions that limit the amount of feedback that can be obtained in the transmitted band are well known.⁶ Broadly speaking, the figure of merit of the electron tubes and the incidental circuit capacitances determine the asymptotic cutoff which limits the amount of feedback that can be obtained in the band. With the flexible repeater circuit, capacitances are rather large because of the severe space restrictions and physical length of the structure. Transit time of 1.8° per megacycle per tube and a like amount for the physical length of the feedback loop reduced the available feedback by 2 db.

Margins of 10 db at phase cross-over and 30° at gain cross-over were set as design objectives. While these may seem to be ultraconservative in view of the tight controls placed on components and the mechanical assembly, it should be borne in mind that the repeaters are inaccessible and repairs would be costly.

Modulation and tube aging considerations require a minimum feed-



- | | |
|------------------------------|------------------------------|
| 1 INPUT TERMINAL | 10 VACUUM TUBE (THIRD STAGE) |
| 2 INPUT BLOCKING CAPACITOR | 11 OUTPUT NETWORK |
| 3 GROUNDING CAPACITOR | 12 BETA NETWORK (1) |
| 4 CRYSTAL | 13 BETA NETWORK (2) |
| 5 INPUT NETWORK | 14 GAS TUBE |
| 6 VACUUM TUBE (FIRST STAGE) | 15 DRYER |
| 7 FIRST INTERSTAGE NETWORK | 16 OUTPUT BLOCKING CAPACITOR |
| 8 VACUUM TUBE (SECOND STAGE) | 17 OUTPUT TERMINAL |
| 9 SECOND INTERSTAGE NETWORK | |

Fig. 9 — Repeater make-up.

back of 33–34 db. With the restrictions noted above and the effect of the potentiometer terms on the available feedback, the top edge of the band is limited to about 165 kc with the desired feedback.

Mechanical Design

To provide a flexible structure the repeater unit is assembled in a number of longitudinal sections mechanically coupled by helical springs and electrically interconnected by means of bus tapes. The assembly is composed of 17 sections. Figs. 9 and 10 show the repeater make-up and an assembled unit.

The sections consist of the circuit component, or components, mounted in machined plastic forms and enclosed in a plastic container which in turn is enclosed in a housing of the same material.* The sections contain circuit components grouped functionally such as input coupling network, interstage, electron tube, or high voltage blocking capacitor. In the case of the feedback network it was necessary to mount the network in two sections because of the large numbers of components involved. A typical network, container and housing are shown in Fig. 11.

The bus tapes are placed in grooves milled in the outer surfaces of the

* The material used is methyl methacrylate which was chosen for its physical and chemical stability and good machinability.

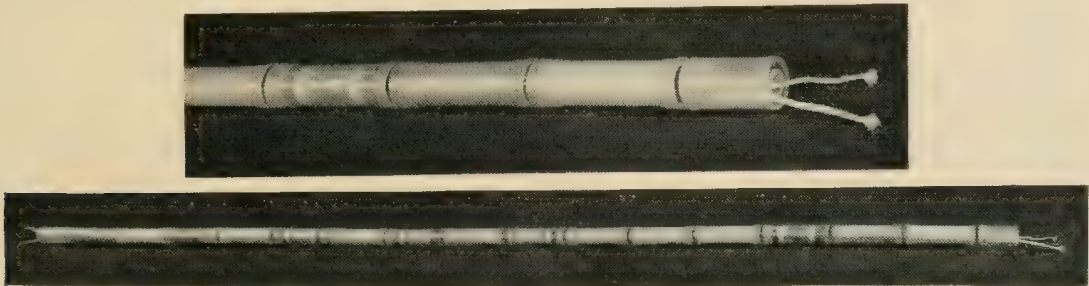


Fig. 10 — Overall view of the repeater unit.

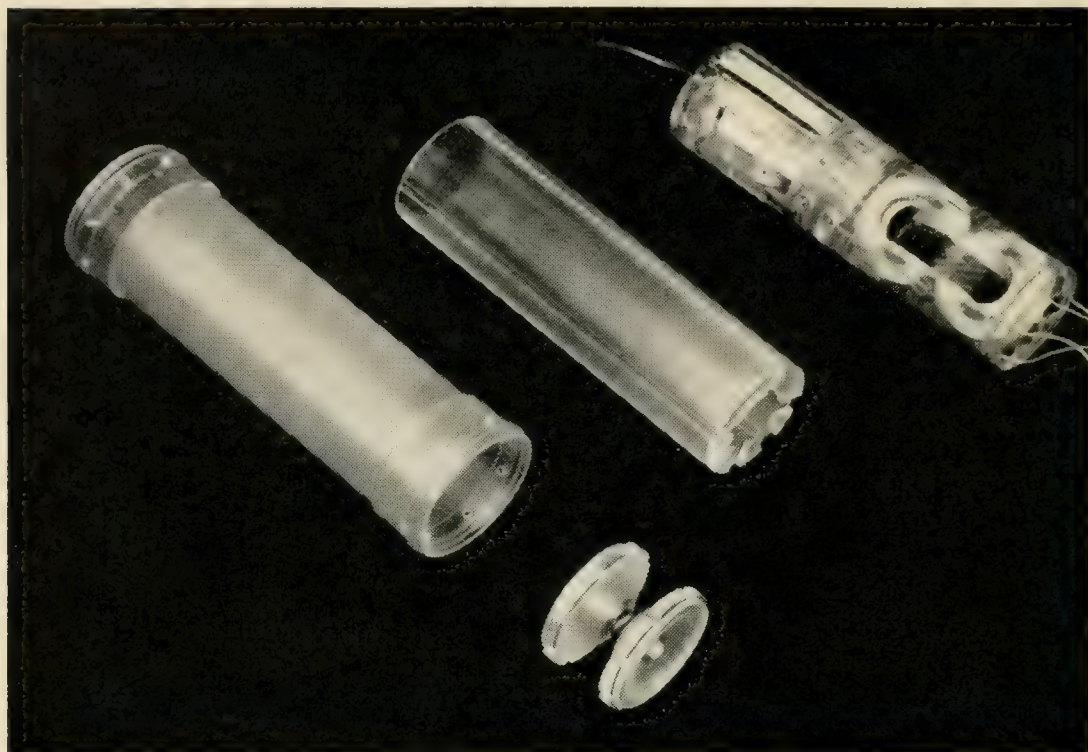


Fig. 11 — Network section.

section containers. Wiring spaces are machined into the ends of the containers for connecting network leads to the bus tapes. The housing is placed over the container and the buses and is closed by a plastic coupler plate which also forms part of the intersection couplers. The coupler plates are fastened to the housing with plastic pins.

Between sections the bus tapes are looped toward the longitudinal axis of the repeater unit. The dimensions of the loop are rigidly controlled so that as the unit is flexed during bending of the repeater, the loops always return to their original location between sections and do not short to each other or the metal outer container. The bus tapes have either an electrical connection or lock at one end of each section to eliminate any tendency of the tapes to creep as the repeater unit is flexed.

The buses consist of two copper tapes in parallel to guard against opens should one tape break. The design of the connections to the buses is such that once the section is closed there can be no disturbance of the tapes or network leads in the vicinity of the electrical connections. The bus-type wiring plan was chosen as the best arrangement for the long structure in keeping with the stringent transmission requirements. Electrically adjacent but physically remote components can thus be inter-

connected with careful control of the parasitic capacitances and couplings to insure reproducibility from unit to unit in manufacture.

COMPONENTS

The development of passive components for use in the flexible repeater presented a number of unusual problems, the most important being: (1) the extreme reliability, (2) the high degree of stability, (3) the limitations on size and shape and, (4) an environment of constant low temperature.

The repeaters for the transatlantic system contain a total of approximately 6,000 resistors, capacitors, inductors and transformers. If we are to be 90 per cent certain of attaining the objective of 20 years service without failure of any of these components, the effective average annual failure rate for the components must be not more than 1 in a million. To assure this degree of reliability by actual tests would require more than 400 years testing on 6,000 components. Obviously some other approach to insure reliability is required. The most obvious avenue, that of providing a large factor of safety, was not open because of space limitations.

Fortunately, with only one exception, the passive components do not wear out. Thus the approach to reliability could be made by one or more of the following:

1. The use of constructions and materials which have been proved by long use, particularly in the Bell System.
2. The use of only mechanically and chemically stable materials.
3. The use of extreme precautions to avoid contamination by materials which might promote deterioration.
4. Special care in manufacture to insure freedom from potentially hazardous defects.

The philosophy of using only tried and proved types of components dictated the use of wire wound resistors, impregnated paper and silvered mica capacitors and permalloy cores for inductors and transformers. While newer and, in some ways, superior materials are known, none of these possessed the necessary long record of trouble-free performance. In some cases, particularly in resistors, this approach resulted in more difficult design problems and also in physically larger components. While the ambient conditions in the repeater, i.e., low temperatures and extreme dryness, are ideal from the standpoint of minimizing corrosion or other harmful effects of a chemical nature, the materials used in the fabrication of components were nevertheless limited to those which are in-

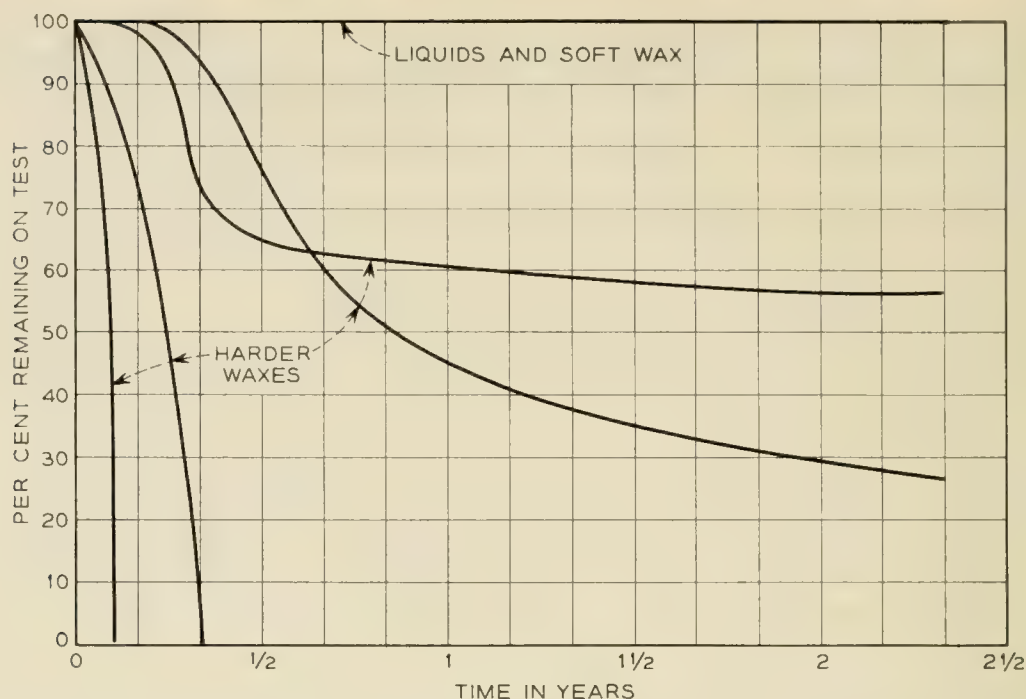


FIG. 12 — Accelerated life tests on paper capacitors with various impregnants at room temperature (60° to 80°F).

herently stable and nonreactive. In addition, raw materials were carefully protected from contamination from the time of their manufacture until they were used, or, wherever possible, they were cleaned and tested for freedom from contaminants just prior to use. Unusually detailed specifications were prepared for all materials.

The effort to achieve extreme reliability also influenced or dictated a number of design factors such as the minimum wire diameters used in wound apparatus, the use of as few electrical joints as possible and the use of relatively simple structures. These limitations resulted in the use of unencased components in most instances. Wherever possible, the ends of windings were used as terminal leads to avoid unnecessary soldered connections. This injected the additional hazard of lead breakage owing to handling during manufacture and inspection. This hazard was minimized in most instances by providing the windings with extra turns which were removed just before the component was assembled in the network. Thus, the lead wires in the final assembly had never been subjected to severe stress. Where this technique was impracticable, special fixtures and handling procedures were used to prevent undue flexing or stressing of lead wires.

As mentioned above there was one type of passive component in which life is a function of time and severity of operating conditions. These are the capacitors, especially those subject to high voltages. Because of this

and the fact that the physical and electrical requirements dictated the use of relatively high dielectric stress in these capacitors, a program of study covering a wide range of dielectric materials was undertaken about 1940. This study showed that none of the usual solid or semisolid materials used to impregnate paper capacitors were suitable for continuous use at sea bottom temperatures. Typical results of this program are shown in Figs. 12 and 13. These curves show the performance of capacitors operating at approximately 1.8 times normal dielectric stress at both sea bottom and room temperatures. It is evident that even semisolid impregnants are inferior to liquids at the lower temperature. The need for the maximum capacitance in a given space restricted the field still further, so that the final choice was a design using castor-oil-impregnated kraft paper as the dielectric.

It is well established that the life of impregnated paper capacitors is inversely proportional to the fourth to sixth power of the voltage stress; or

$$\frac{L_1}{L_2} = \left(\frac{V_2}{V_1} \right)^p$$

where p ranges from 4 to 6. This fact permits the accumulation of a large amount of life information in a relatively short time. In order to insure

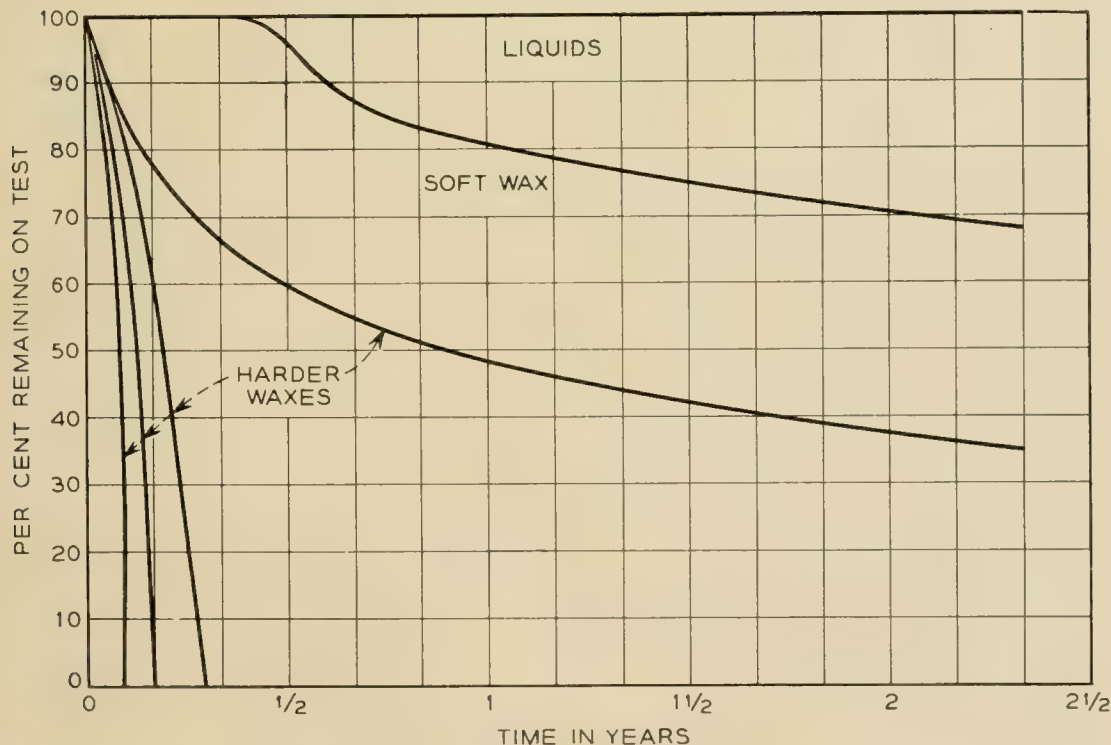


Fig. 13 — Accelerated life tests on paper capacitors with various impregnants at 40°F.

that the capacitor design selected would provide the degree of reliability required, a number of capacitors were constructed and placed on test at voltage stresses ranging from $1\frac{1}{2}$ to $2\frac{1}{4}$ times the maximum stress expected in service. From the performance of these samples, a prediction of performance under service conditions can be made as follows:

The total equivalent exposure in terms of capacitor years at the maximum service voltage can be computed for the samples under test by the following summation:

$$T = N_1 T_1 \left(\frac{V_1}{V_s} \right)^p + N_2 T_2 \left(\frac{V_2}{V_s} \right)^p + \cdots + N_r T_r \left(\frac{V_r}{V_s} \right)^p \quad (1)$$

where $N_1, N_2, \cdots N_r$ are the number of samples on test at voltage stresses V_1, V_2 and V_r , $T_1, T_2, \cdots T_r$ are the total times of the individual tests and V_s is the maximum voltage stress under service conditions.

If, as has been the case in the tests described above, there has been only one failure in the total exposure T , we can estimate from probability equations the limits or bounds within which the first failure will occur in a system involving a given number of capacitors operating at a voltage stress V_s . These equations are:

$$\text{probability of no failures in exposure } T = e^{-(T/L)} \quad (2)$$

probability of more than one failure in exposure time $T =$

$$1 - \left(1 + \frac{T}{L} \right) e^{-T/L} \quad (3)$$

where T is obtained from (1) and L is the total exposure in the same units as T for the service conditions. The solutions of (2) and (3) for L using any desired probability give the maximum and minimum exposures in capacitor-years, within which the first failure may be expected to occur under service conditions.

However, since the voltage on the capacitors varies from repeater to repeater, it is necessary to determine the equivalent exposure of the system in terms of capacitor-years per year of operation at the maximum service voltage in order to estimate the time to the first failure in the system. This is obtained from (1) for one-half of one cable by substituting the supply voltage at each repeater for V_1, V_2 , etc., the maximum service voltage for V_s and the number of capacitors per repeater for N . The total exposure for a two cable system is then 4 times this figure. With the data which has been accumulated and the number of capacitors and voltages of the transatlantic system, we estimate with a probability of being correct nine times in ten that the first "wear-out" failure of a

capacitor in the transatlantic system will not occur in less than 16 years nor more than 600 years.

There is, of course, the possibility of a catastrophic or early failure due to mechanical or other defects not associated with normal deterioration of the dielectric. Such potential failures are not always detected by the commonly used short-time over-voltage test. Thus, for submarine cable repeaters, all capacitors subjected to dc potentials in service are subjected to at least $1\frac{1}{2}$ times the maximum operating voltage for a period of four to six months before they are used in repeaters. Experience indicates that this is adequate to detect potential early failures. The results of this type of testing on submarine cable capacitors is an indication of the care used in selecting materials and manufacturing the capacitors. Only one failure has occurred in more than 3,000 capacitor-years of testing.

An important aspect of the control of quality of components is the control of the raw materials used in their manufacture. For the transatlantic project, this was accomplished by rigid specifications, thorough inspection and testing, supplemented in some cases by a process of selection.

This can be illustrated by the procedure used for selecting the paper used as the dielectric in capacitors. The Western Electric Company normally inspects many lots of capacitor paper during each year. Those lots which were outstanding in their ability to stand up under a highly accelerated voltage test were selected from this regular inspection process. These selected lots were then subjected to a somewhat less highly accelerated life test. Paper which met the performance requirements of this test was slit into the proper widths for use in capacitors. Sample capacitors were then prepared with this paper and so selected that they represented a uniform sampling of the lot at the rate of one sample for approximately each three pounds of paper. These samples were impregnated with the same lot of oil to be used in the final product. Satisfactory completion of accelerated life and other tests on these samples constituted final qualification of the paper for production of capacitors. Relatively few raw materials were adaptable to such tests or required such detailed and exhaustive inspection as capacitor paper. But the attitude in all cases was that the material be qualified not only as to its primary constituents or characteristics but also as to its uniformity and freedom from unwanted properties.

To a considerable extent, stability of components is assured by the practice of using only those types of structures which have long records of satisfactory field performance. However, in some cases, a product far

more stable than usual was required. This was true of the high voltage capacitors where other requirements dictated the use of impregnated paper as the dielectric but where the degree of stability required was comparable to that expected of more stable types of capacitors. In so far as possible, stability was built into the components by appropriate design but, where necessary, stabilizing treatments consisting of repeated temperature cycles were used to accelerate aging processes to reach a stable condition prior to assembly of the repeaters. Temperature cycling or observation over periods up to six months were used also to determine that the components' characteristics were stable.

Exceptional inspection procedures followed to insure reliability and stability are described in detail in a companion paper.⁷

As mentioned earlier, the design and construction of components was simplified by omitting housings or containers, except for oil impregnated paper capacitors. Adequate mountings for the components were obtained in several ways. Mica capacitors were cemented to small bases of methyl methacrylate which were in turn cemented in suitable recesses in network structures. Inductors and transformers were cemented directly into recesses in the network housings. Fig. 14 illustrates some of these structures and their mounting arrangements. On the bottom is a molybdenum permalloy dust core coil in which a mounting ring of methyl

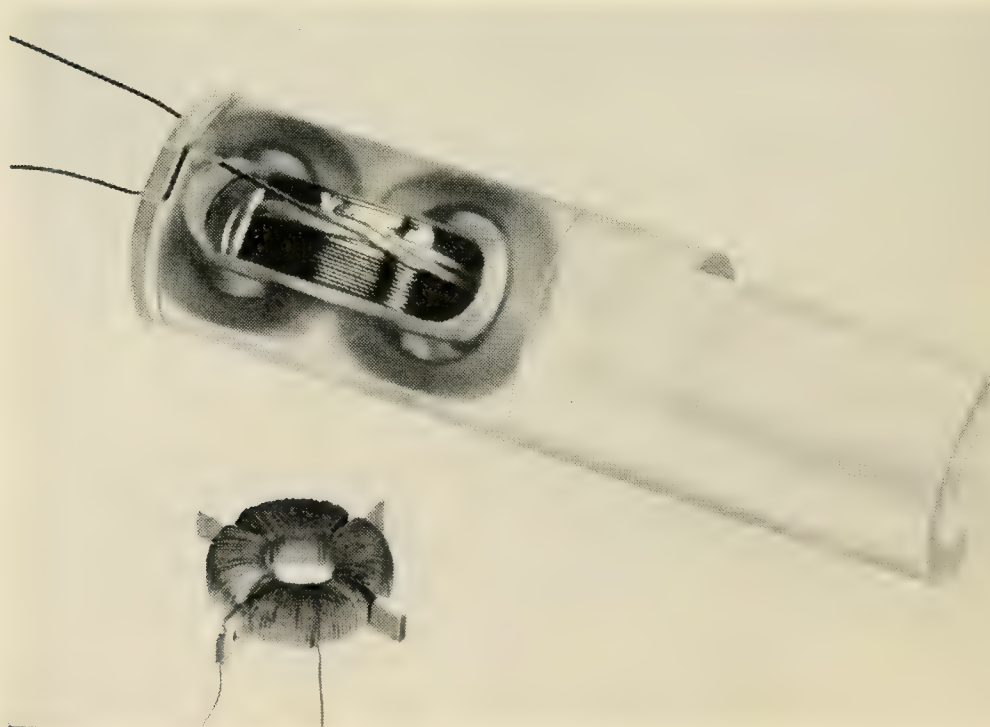


Fig. 14 — Mounting for molybdenum permalloy dust core coils.

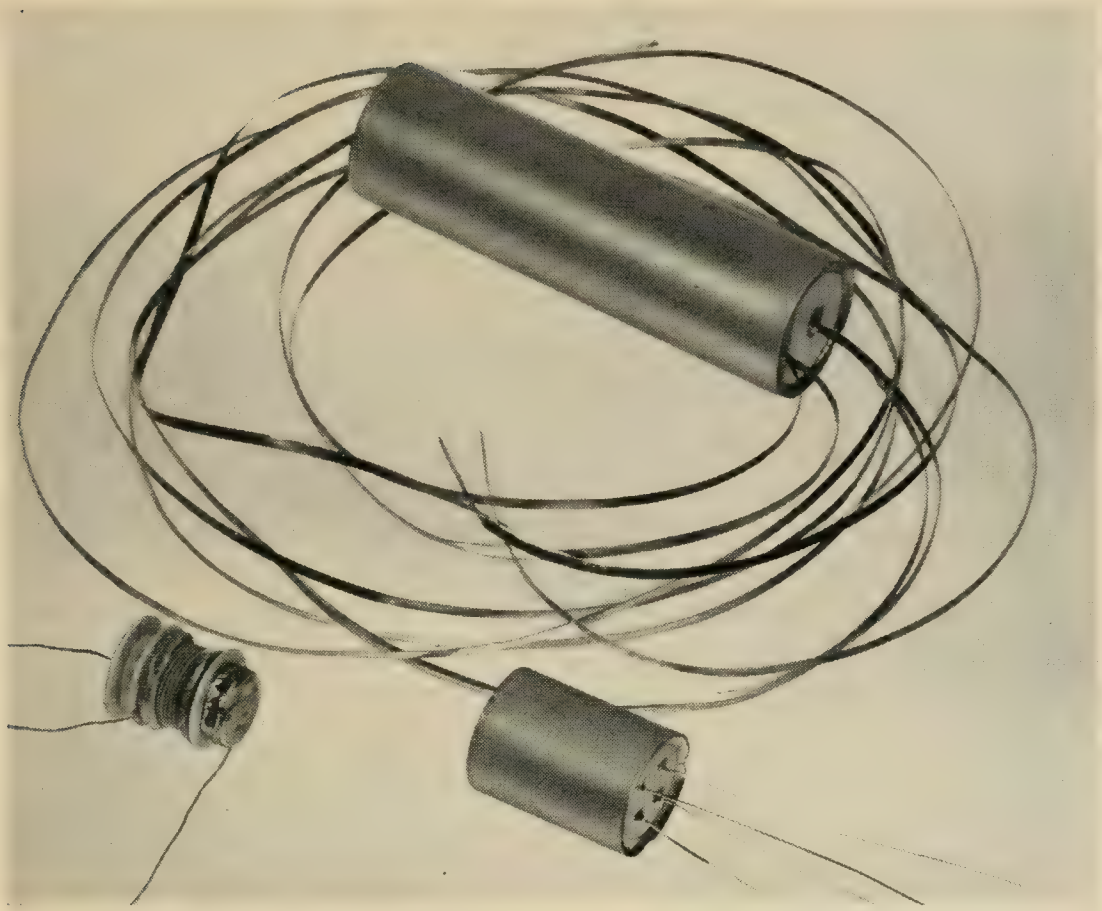


Fig. 15 — Capacitor and resistor capacitor combinations.

methacrylate provided with radial fins is secured around the core by tape and the wire of the winding. Such inductors were mounted by cementing the projecting fins into slots arranged around a recess in the network housing. On top is an inductor which, for electrical reasons, required a core of greater cross-section than could be accommodated in the network when made by the usual toroidal construction. In this case, the effective cross-section of two cores was obtained by cementing the cores in a "figure-8" position and by applying the winding so that it threads the hole in both cores. With these constructions, the cement used to secure the inductors does not come into contact with the wire of the winding which is thereby not subject to strains produced by curing of the cement.

For economy of space and also to reduce the number of soldered connections, many of the components' structures contain two or more elements. Inductors and resistors were combined by winding inductors with resistance wire. Separate adjustment of inductance and resistance were obtained by adjusting turns for inductance and the length of wire in a

small "non-inductive" winding for resistance. The inductor on the bottom in Fig. 14 illustrates one type in which the non-inductive part of the winding is placed on one of the separating fins. In some cases, capacitors and resistors were also combined. Fig. 15 shows two of these. The capacitor at the bottom right contains three capacitances and a single resistance in the same container. This construction requires that the resistor parts be capable of withstanding the capacitor drying and impregnation process and also that the resistor contain nothing which would be harmful to the capacitor. The capacitor on the left in this figure is housed in a ceramic container on which is wound a resistor. The capacitor at the top is a high-voltage type which, aside from electron tubes, represents the largest single component used in the repeater. In this capacitor, the tape terminals which contact the electrodes are brought out through the ceramic cover and are made long enough to reach an appropriate point so as to avoid additional soldered connections. Such special designs introduced many problems in the manufacture of the components. However, the improved performance of the repeater and the increase in the inherent reliability of the overall system fully justified the greater effort which was required for the production of such specialized apparatus.

POWER BY-PASS GAS TUBE*

The fault locating means, referred to previously, requires that the power circuit through the cable be continuous. To protect against an open circuit in the repeater, such as a heater failure, an additional device is required to bypass the line current. This bypass must be a high resistance under normal operating conditions since any current taken by this device must be supplied through preceding repeaters. If an open circuit occurs the bypass must carry the full cable current. At full current, the voltage drop should be small to avoid excessive localized power dissipation in the repeater. The device should recover when power is removed so that false operation by a transient condition will not permanently bypass the repeater.

A gas diode using an ionically heated cathode has been used to meet these requirements. By making the breakdown voltage safely greater than the drop across the heater string, no power is taken by the tube under normal repeater operation. In the event of an open circuit in the repeater, the voltage across the tube rises and breakdown occurs. Full cable current is then passed through the gas discharge. Removal of power

* Material contributed by Mr. M. A. Townsend.

from the cable allows the tube to deionize and recover in the event of false triggering by transients. The cathode is a coil of tungsten wire coated with a mixture of barium and strontium oxide. A cold cathode glow discharge forms when the tube is first broken down. This discharge has a sustaining voltage of the order of 70 volts. The glow discharge initially covers the entire cathode area. Local heating occurs and some parts of the oxide coating begin to emit electrons thermionically. This local emission causes increased current density and further increases the local heating. The discharge thus concentrates to a thermionic arc covering only a portion of the coil. The sustaining voltage is then of the order of 10 volts.

Mechanically the tube was designed to minimize the possibility of a short circuit resulting from structural failure of tube parts. Fig. 16 shows the construction of the tube. The glass envelope and stem structure which had previously been developed for the hot cathode repeater tubes were used as a starting point for the design. The anode is a circular disk of nickel attached to two of the stem lead wires. To provide shock resistance the supporting stem leads are crossed and welded in the center. To protect against weld failure, a nickel sleeve is used at each end of the cathode

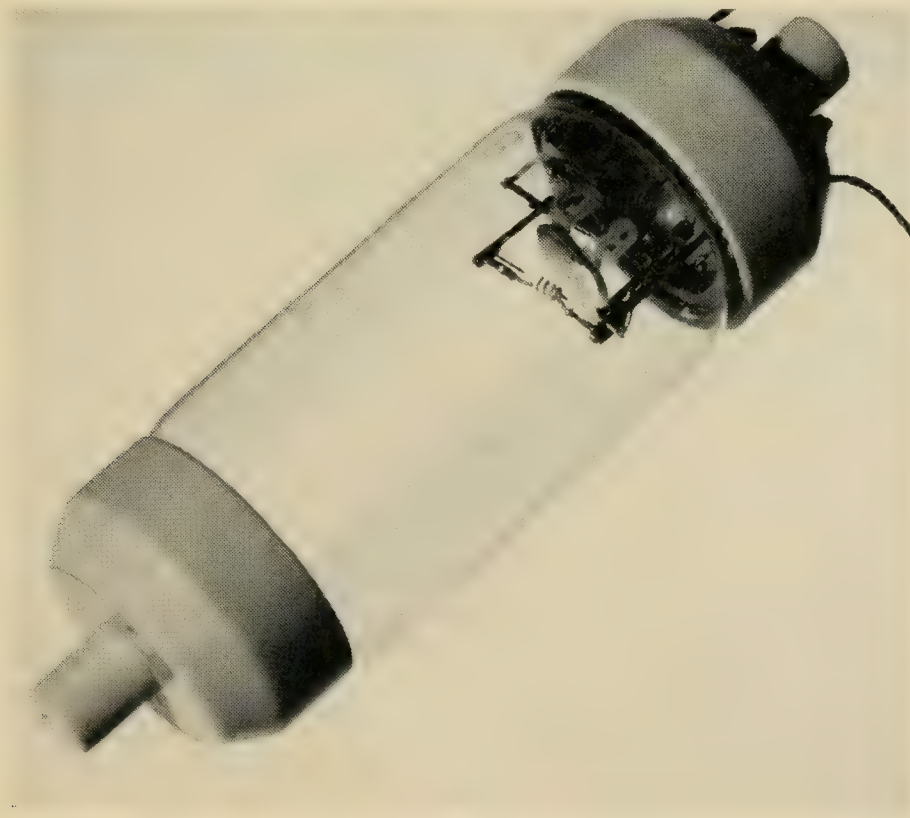


Fig. 16— The power by-pass gas tube.

coil. It is crimped to hold the coil mechanically in place and then welded at the end for electrical connection. At the end of the coil as well as in all other places where it is possible, a mechanical wrap is made in addition to spot welding. An additional precaution is taken by inserting an insulated molybdenum support rod through the center of the cathode coil. The filling gas is argon at a pressure of 10 mm Hg. To provide initial ionization, 1 microgram of radium in the form of radium bromide was placed on the inside of the tube envelope. All materials were procured in batches of sufficient size to make the entire lot of tubes and carefully tested before being approved for use. The tubes were fabricated in small groups and a complete history was kept of the processing of each lot.

For detailed study of tube performance, a number of electrical tests were made. These involved measurements of breakdown voltage, operating voltage as a glow discharge at low current, current required to cause the transition to a thermionic arc, the time required at the cable current to cause transition to the low voltage arc, and the sustaining voltage at the full cable current.

All tubes were aged by operating at 250 milliamperes on a schedule which included a sequence of short on-off periods (2 min. on, 2 min. off) followed by periods of continuous operation. A total of 150 starts and 300 hours of continuous operation were used. Following this aging schedule the tubes were allowed to stabilize for a few days and then subjected to a 2-hour thermal treatment or pulse at 125°C. It was required that no more than a few volts change in breakdown voltage occur during this thermal pulse before a tube was considered as a candidate for use in repeaters.

After aging and selection as candidates for repeaters, tubes were stored in a light-tight can at 0°C. Measurements were made to assure stability of breakdown voltage and breakdown time.

The quality of each group of 12 tubes was further checked by continuous and on-off cycling life tests. The fact that none of these tubes has failed on the cycling tests at less than 3,500 hours and 1,500 starts and no tube on continuous operation has failed at less than 4,200 hours gives assurance that system tubes will start once and operate for the few hours necessary to locate a defective repeater. Long-term shelf tests of representative samples at 70°C and at 0°C give assurance of satisfactory behavior in the system.

CONTAINER AND SEALS

The design of the flexible enclosure for the flexible repeater unit is basically the same as it emerged from its development stages in the

1930's. It is virtually identical to the structure of the repeaters manufactured by the Bell Telephone Laboratories for the cables laid in 1950 between Key West and Havana.³

The functions of the enclosure are to protect the repeater unit from the effects of water at great pressure at the ocean bottom; to provide means of connecting the repeater to the cable before laying; and to be slender and flexible enough to behave like cable during laying. How these functions are met in the design may be more readily understood by reference to Fig. 17.

The repeater unit, described earlier, is surrounded by a two-layer carcass of steel rings, end to end. The rings are surrounded in turn by a copper tube $1\frac{3}{4}$ inches in diameter and having a $\frac{1}{32}$ -inch wall.

When a repeater is bent during laying by passing onto the cable-ship drum, the steel rings separate at the outer periphery of the bend and the copper tube stretches beyond its elastic limit. As the repeater leaves the drum under tension the rings separate and the copper stretches on the opposite side, leaving the repeater in a slightly elongated state. At the ocean bottom, hydraulic pressure restores the repeater to its original condition with rings abutted and the copper tube reformed.

The system of seals in each end of the tube consists of (1) a glass-to-Kovar seal adjacent to the repeater unit, (2) a rubber-to-brass seal seaward from the glass seal, and (3) a core tube and core sleeve seal seaward from the rubber seal.

The glass seal, although capable of withstanding sea bottom pressures, is primarily a water vapor barrier and a lead-through for electrical connection to the repeater circuit. In service it is normally protected from exposure to sea pressures by the rubber seal.

The rubber seal, capable of withstanding sea bottom pressures, is indeed exposed to these pressures for the life of the repeater, but is not exposed to sea water. It is likewise a lead-through for electrical connection from the cable to the glass seal.

The core sleeve seal is an elastic barrier between sea water on the outside and a fluid on the inside. This fluid, polyisobutylene, is a viscous honey-like substance, chemically inert, electrically a good insulator, and a moderately good water vapor barrier. It fills the long thin annular space outside the cable core and inside a copper core tube and thus becomes the medium of transmitting to the rubber seal the sea pressure exerted on the core sleeve. It can be seen that the core sleeve seal has nominally no pressure resisting function and no electrical function.

The same fluid is also used to fill the space between the glass and rubber seals. Voids at any point in the system of seals are potential hazards to long, trouble-free life. Empty pockets, for instance, lying between the

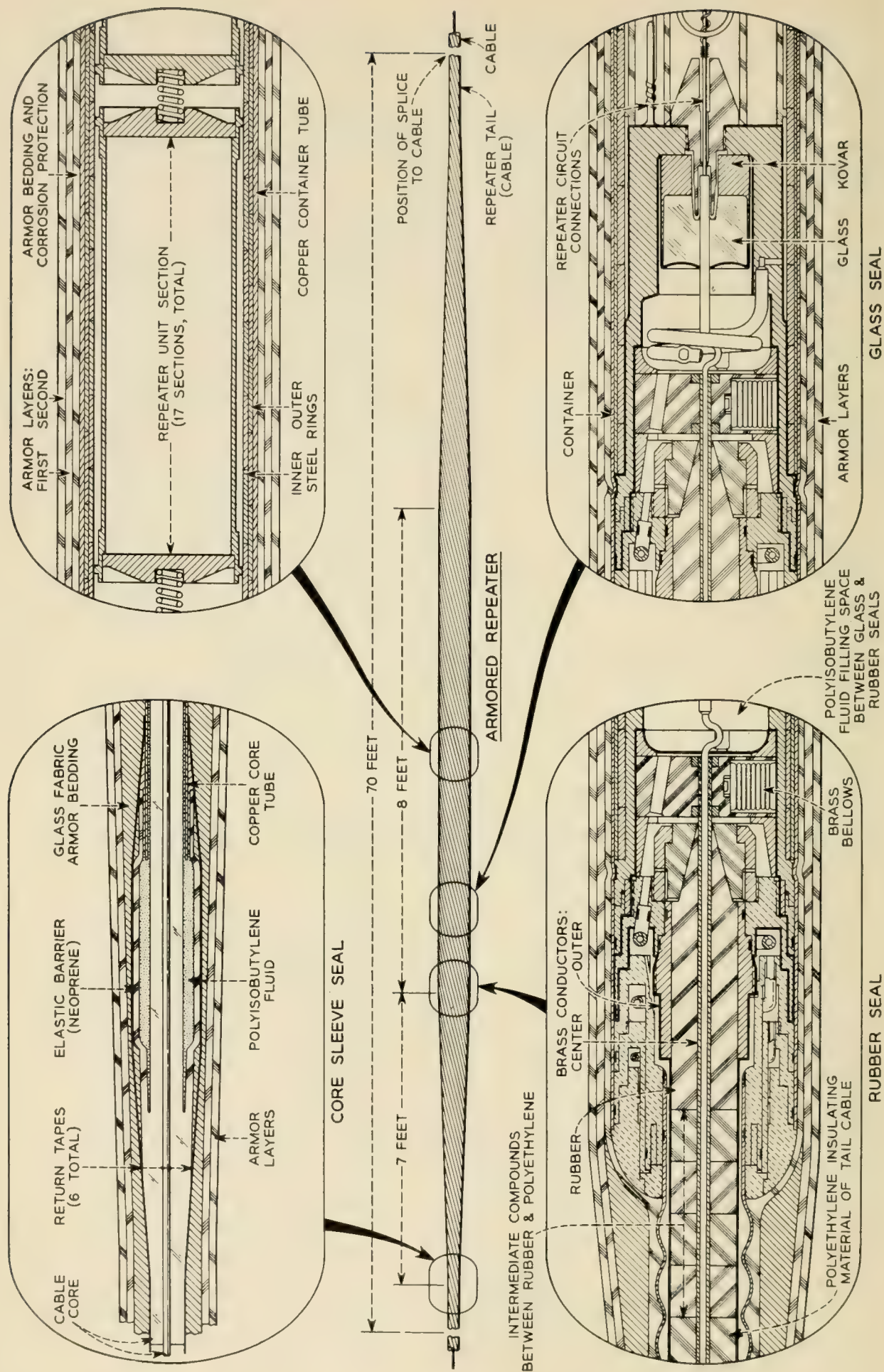


Fig. 17 — Details of container and seals.

central conductor and the outer conductor, or container, are capable of becoming electrically conducting paths if filled with water vapor. As pointed out in companion papers,^{2, 4} the voltage between the repeater (and cable) central and outer conductors is in the neighborhood of 2,000 volts at the ends of the transatlantic system.

The filling of the seal interspace with a liquid would defeat one function of the rubber seal if special features were not provided in the rubber seal design. Very slight displacement of the rubber seal toward the glass seal because of sea pressure, or resulting from reduction in volume owing to falling temperature, would otherwise build up pressure in the liquid and on the glass seal. We avoid this by providing a kind of resilience in the interspace chamber. Three small brass bellows, partly compressed, occupy fixed cavities in the chamber. They can compress readily and maintain essentially constant conditions independent of external pressures and temperatures.

The entire repeater assembly enclosed in copper is approximately 23 feet long. Tails of cable at each end make the total length about 80 feet before splicing. The central conductor of each cable tail is joined to the rubber seal central conductor, with the insulation molded in place in generally the same manner as in cable-to-cable junctions elsewhere in the system. The outer-conductor copper tapes of the cable tails are electrically connected to the copper core tubes.

The copper region is coated with asphalt varnish and gutta percha tape to minimize corrosion. Over this coating bandage-like layers of glass fabric tape are built up to produce an outer contour tapering from cable diameter at one end up to repeater diameter and back down to cable diameter at the opposite end. The tape covering is saturated with asphalt varnish. This tape is primarily a bedding for the armor wires that are laid on the outside of both cable tails and repeater to make the repeater cable-like in its tensile properties and capable of being spliced to cable.

In the region of the repeater proper where the diameter is double that of cable, extra armor wires are added to produce a layer without spaces. Also, to avoid subjecting the repeater to the torque characteristically present in cable under the tensions of laying, a second layer of armor wires of opposite lay is added over the first layer. This armoring process is so closely related to the armoring of cable core in a cable factory that it is performed there.

Materials

Following the same design philosophy applied to the repeater components, the materials of construction of the repeater container and seals

were chosen for maximum life, compatibility with each other, and for best adaptability to the design intent. Specifications particularly adapted to this use were set up for all of the some 50 different metals and non-metals employed in the enclosure design. In general, the methods established for proving the integrity of the materials are more elaborate than usual commercial practice. In most instances, such as that of copper container tubes, the extraordinary inspection for defects and weaknesses with its resulting rejection rate, resulted in high cost for the usable material.

TESTING

A substantial part of the development work on the repeater enclosure was concerned with devising tests that give real assurance of soundness and stability. It is beyond the scope of this paper to discuss how each part is tested before and after it is assembled but certain outstanding tests deserve mention.

Steel Ring Tests

Each of the inner steel rings, before installation, is required to pass a magnetic particle test to find evidence of hidden metallurgical faults. Each ring is later a participant in a group test under hydraulic pressure simulating the crushing effect of ocean bottom service but exceeding the working pressures. The magnetic particle test is repeated.

Helium Leak Tests

Both glass and rubber seal assemblies, before being installed in repeaters, are required to undergo individual tests under high-pressure helium gas. Helium is used not only because its small molecules can pass through smaller leaks than can water molecules but because of the excellent mass spectrometer type of leak detectors commercially available for this technique. While helium is applied at high pressure to the outer wall of the seal, the inner wall is maintained under vacuum in a chamber joined with the leak detector. The passage of helium through a faulty seal at the rate of 10^{-9} milliliters per second can be detected. Stated differently, this is 1 milliliter of helium in 30 years. The relation of water-leak rate to helium-leak rate is dependent on the physical nature of the leak, but if they were assumed to be equal rates, the amount of water which might enter a tested repeater in 20 years would be 0.66 grams. A desiccant within the repeater cavity is designed to keep the

relative humidity under 10 per cent if the water intake were five times this amount.

After glass seals are silver brazed into the ends of the copper tube of the repeater the helium test is repeated to check the braze and to re-check the seal. For this test the entire repeater must necessarily be submerged in high pressure helium. Obviously, in order to sense a possible passage of the gas from the outside to the inside, the leak detector vacuum system must be connected to the internal volume of the repeater. For this and other reasons a small diameter tube that by-passes the seal is provided as a feature of the seal design. After the leak integrity of the repeater is established by this means for all but the access tube, this tube is then used as a means of vacuum drying the repeater and then filling it with extremely dry nitrogen. Following this, the tube is closed by welding and brazing. This closure is then the only remaining leak possibility and is checked by a radioisotope leak test.

Radioisotope Leak Test

Of various methods of detecting the passage of very small amounts of a liquid or a gas from the outside to the inside of a sealed repeater, a scheme using a gamma-emitting radioisotope appeared to be the most applicable.

The relatively small region of the welded tube referred to above is surrounded by a solution of a soluble salt of cesium¹³⁴. With the entire repeater in a pressure tank, hydraulic pressure in excess of service pressures is applied for about 60 hours. The repeater is removed from the tank, the radioactive solution is removed and the test region is washed by a special process so as to be essentially free from external radioactivity. A special geiger counter is applied to the region. If there has been no leak the gamma radiation reads a low value. If an intake has occurred of as much as one milligram of the isotope solution, the radiation count is about four to five times greater than that of the no-leak condition. The rate of leak indicated is an acceptable measure of soundness of the repeater closure.

The helium and subsequent isotope leak tests are made on a repeater not only when its glass seals are installed but are performed again on each rubber seal after it is brazed in place.

Electrical Tests

Prior to assembly into the repeater the various networks are tested under conditions simulating as nearly as is feasible the actual operating

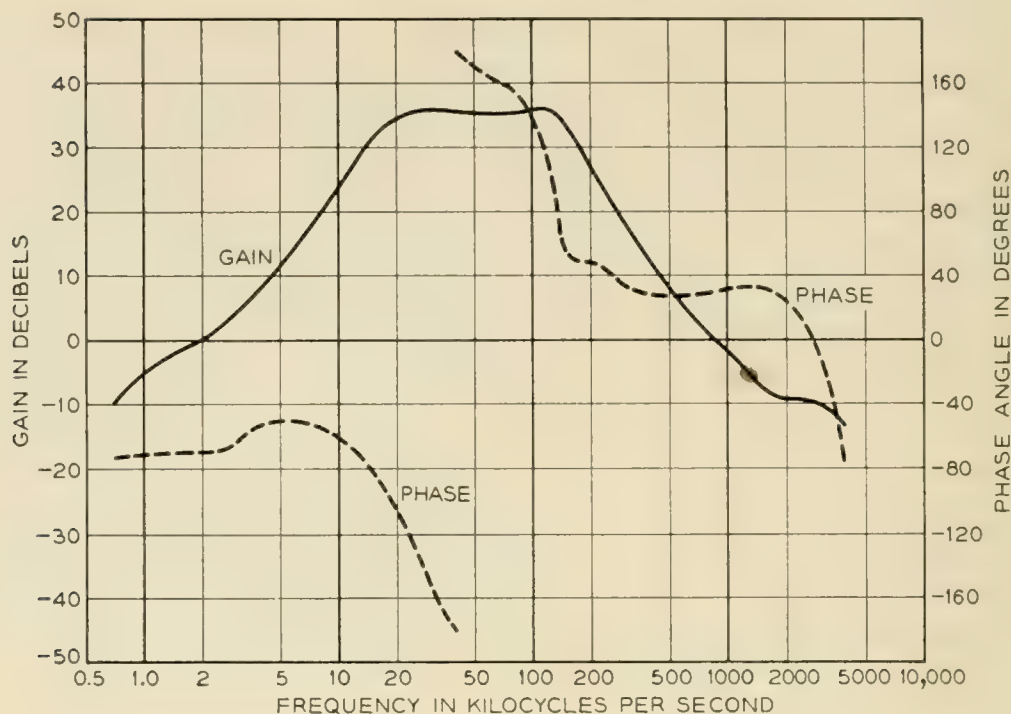


Fig. 18 — Mu Beta gain and phase.

conditions of the particular network. The input and output coupling networks and the beta networks enter directly into the insertion gain and hence are held to very close limits. To ensure meeting these limits elements which go into a particular network are matched and adjusted as a group before assembly into the network.

Repeater units are tested for transmission performance both before and after closing. These tests consist of; mu-beta measurements (simultaneous measurements of gain and phase of the feedback loop); noise; modulation; insertion gain at many frequencies; exact frequency of the fault location crystal and crystal peak gain. Modulation and crystal frequency measurements are made with the repeater energized at 225 milliamperes cable current and also at 245 milliamperes as a check on the ultimate performance of the whole system initially and after aging.

PERFORMANCE OF REPEATERS

The phase and gain characteristics of the feedback loop of the repeater are shown in Fig. 18. It will be noted that at the upper edge of the band the feedback is a little less than the 33-34 db set as the objective. Additional elements could have been used in the interstages to increase the feedback but the return per element is small. Since any element is a potential hazard, the lower feedback is acceptable.

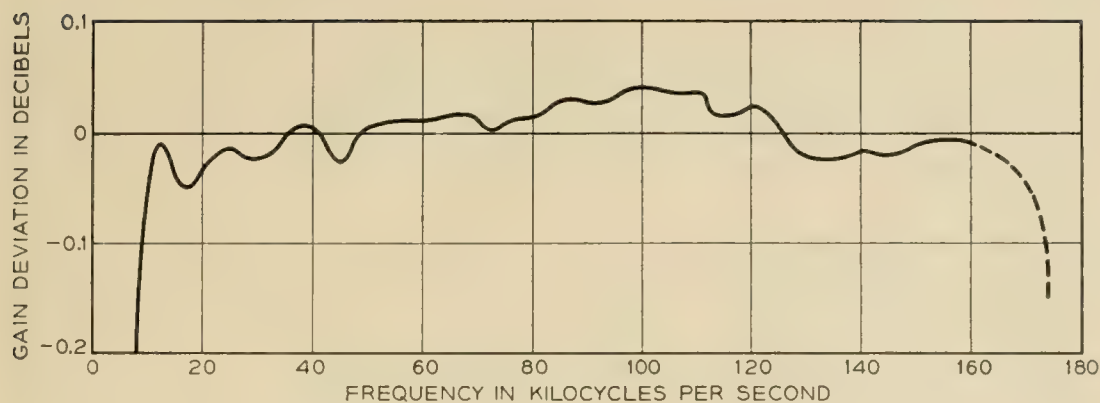


Fig. 19 — Repeater deviation from 36.9 NM design cable.

The deviation of the insertion gain of the repeater from the loss of 36.9 nautical miles of design cable⁵ at sea bottom is shown in Fig. 19. This is well within the objective of ± 0.05 db.

It has been pointed out that the repeater input and output impedance do not match the cable impedance. This results in ripples in the system frequency characteristic due to reflections at the repeater. These are shown in Fig. 20.

The noise performance of the repeater is determined by the input tube and the voltage ratio of the input coupling network. Amplifier noise referred to the input is shown in Fig. 21. At the upper frequencies the

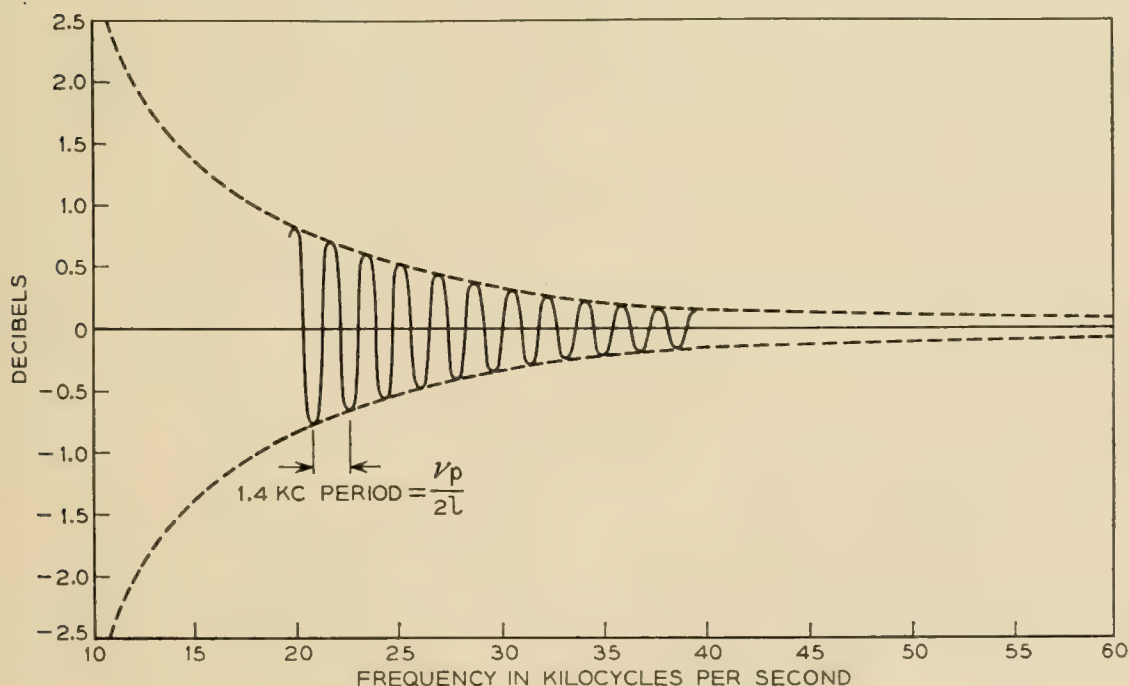


Fig. 20 — Interaction ripple for TAC system.

repeater contribution to cable noise is very small. At the lower frequencies, while the repeater noise is considerably greater than thermal noise, this does not degrade performance because of the lower cable attenuation at these frequencies.

MANUFACTURING DRAWINGS

Because of the extraordinary nature of many of the manufacturing problems associated with undersea repeaters it was determined at the outset that a so-called single-drawing system would be used. For this reason, considerably more information is supplied than is normal. The effect is illustrated best in the rather large number of drawings that consist of text material outlining in detail a specific manufacturing technique. Such drawings specify the devices, supplies and work materials needed to perform an operation, and the step-by-step procedure. Of course, these papers are by no means a substitute for manufacturing skill. Primarily they insure the continuance of practices proved to be effective with the Havana-Key West project.

REPAIR REPEATER

The "repair repeater," used to offset the attenuation of the excess cable which must be added in making a repair, is basically the same general design as the line repeater. It employs a two-stage amplifier, designed to match the loss of 5.3 nautical miles of cable to within ± 0.25 db. The larger deviation compared to the line repeater is permissible since few repair repeaters are expected to be added in a cable. The input

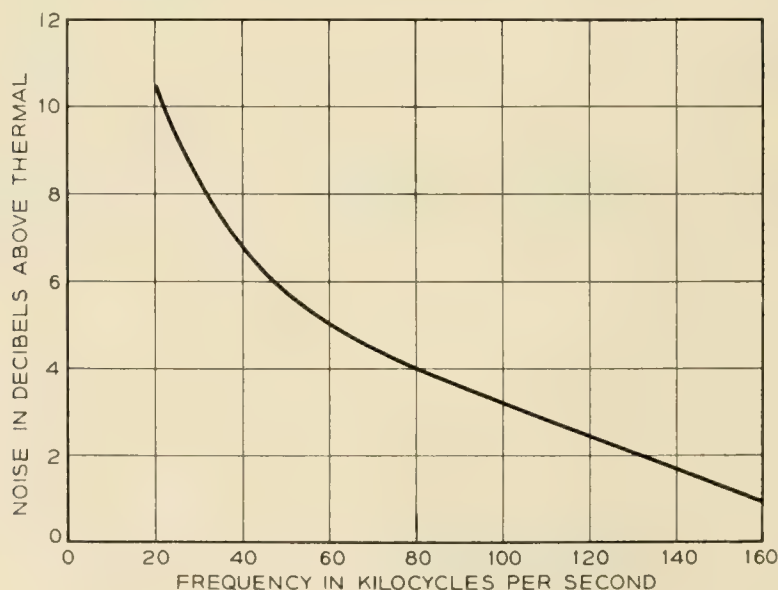


Fig. 21 — Repeater noise.

and output impedances match the cable. As in regular repeaters a crystal and gas tube are provided for maintenance testing. The crystals give approximately 25 db increase in gain and are placed between 173.5 and 174.1 kc so as not to duplicate any frequencies used in the line repeaters. The crystal frequency spacing is 100 cycles.

Wherever possible the same components and mechanical details are used in the repair repeaters as in the line repeaters. When changes in design were necessary, these were modifications in the existing designs rather than new types. Capacitors are like those of line repeaters. Except for the length of the container, the enclosure is identical to the line repeater.

Noise and overload considerations restrict the location of a repair repeater to the middle third of a repeater section.

UNDERSEA EQUALIZERS

Even though the insertion gain of the line repeater matches the normal loss characteristic of the cable rather closely, uncertainties in the knowledge of the attenuation of the laid cable can lead to misalignment which, if uncorrected, would seriously affect the performance of the system. Misalignment which has cable loss shape can be corrected by shortening or lengthening the cable between repeaters at intervals as the cable is laid. Other shapes, however, require the addition of networks or equalizers in the line.

With these factors in mind a series of undersea equalizers were designed. The loss shapes were chosen on the basis of a power series analysis of expected misalignments. The designs were restricted to series im-

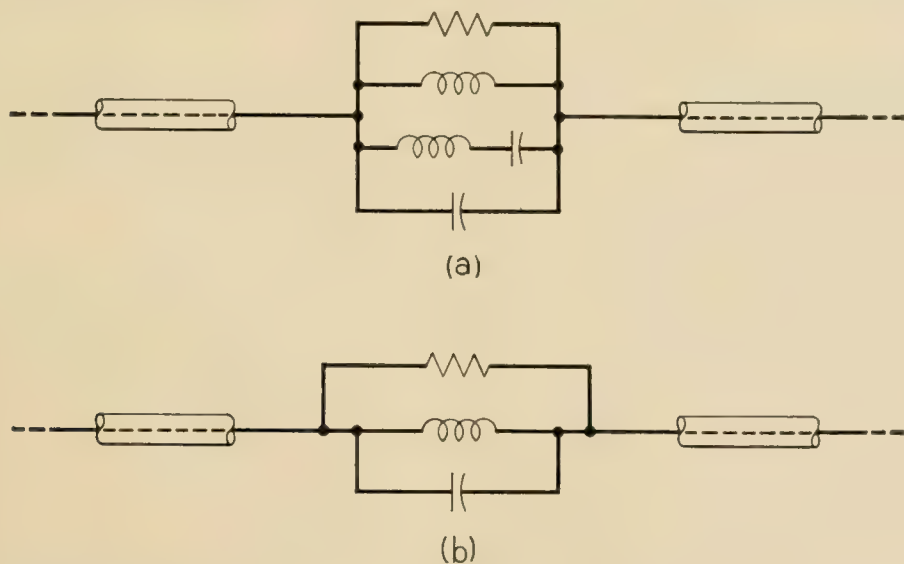


Fig. 22 — (a) Schematic of Type IV equalizer. (b) Schematic of Type V equalizer

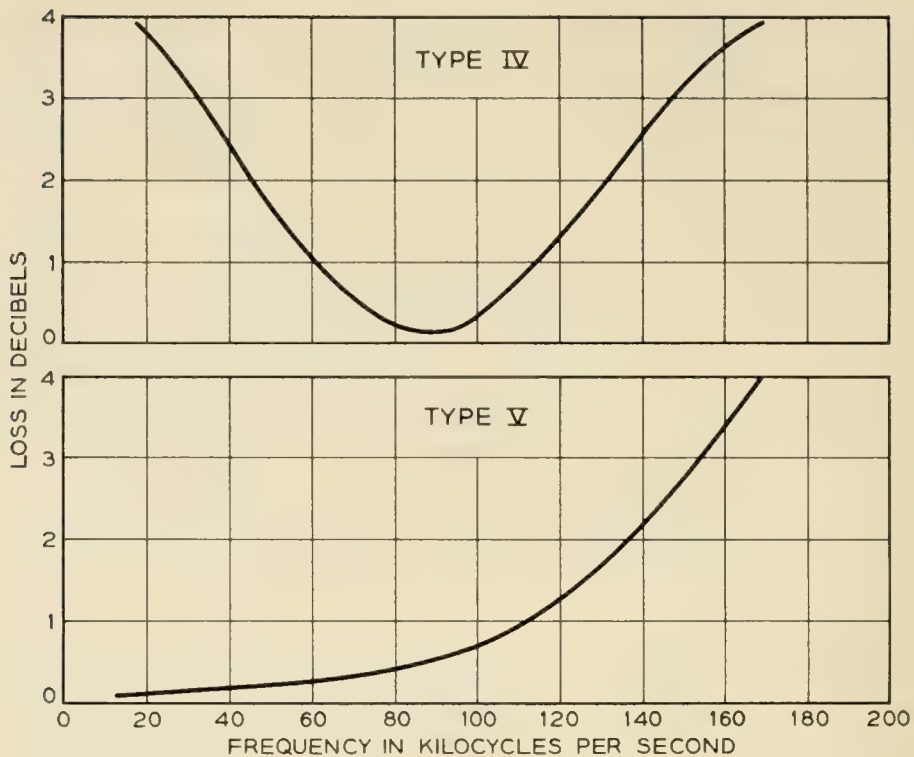


Fig. 23 — Equalizer loss characteristics.

pedance type equalizers to avoid the necessity for shunt arms and the accompanying high-voltage blocking capacitor required to isolate the cable power circuits. This restriction confines the ultimate location of the equalizers to the middle portion of repeater sections to minimize the reaction of the poor repeater impedance on the equalizer characteristic. The dc resistance of equalizers is low so that material increase of the system power supply voltage is not required.

The configuration of two of the equalizers are shown in Fig. 22. The loss characteristics are shown in Fig. 23. Each equalizer has a maximum loss spread in the pass band of about 4 db which represents a compromise between keeping the number of equalizers low and at the same time keeping the misalignment within tolerable limits.

The components used are modifications of the repeater components. The mechanical construction is identical to the repeater except that with the smaller number of elements, the container is materially shorter than a repeater.

ACKNOWLEDGMENTS

Scores of individuals have contributed to the development of these repeaters, some leading to basic decisions, some creating, adapting and

perfecting both electrical and mechanical designs. Many of these people have furnished the continuing drive and enthusiasm that are so essential for a team of engineers and scientists having divergent interests. It is nearly impossible to assign relative importance to the work of transmission engineers, apparatus designers, mathematicians and research scientists in the fields of materials and processes. Equally difficult is any realistic appraisal of the work of all of the technical aides and shop personnel whose contributions are so significant to the final product. The authors of this paper, in reporting the results, therefore acknowledge this large volume of effort without listing the many individuals by name.

REFERENCES

1. E. T. Mottram, R. J. Halsey, J. W. Emling and R. G. Griffith, Transatlantic Telephone Cable System — Planning and Over-All Performance. Page 7 of this issue.
2. H. A. Lewis, R. S. Tucker, G. H. Lovell and J. M. Fraser, System Design for the North Atlantic Link. See page 29 of this issue.
3. J. J. Gilbert, A Submarine Telephone Cable with Submerged Repeaters, *B.S.T.J.*, **30**, p. 65, 1951.
4. G. W. Meszaros and H. H. Spencer, Power Feed Equipment for the North Atlantic Link. See page 139 of this issue.
5. A. W. Lebert, H. B. Fischer and M. C. Biskeborn, Cable Design and Manufacture for the Transatlantic Submarine Cable System. See page 189 of this issue.
6. H. W. Bode, Network Analysis and Feedback Amplifier Design, D. Van Nostrand Co., Inc.
7. H. A. Lamb and W. W. Heffner, Repeater Production for the North Atlantic Link. See page 103 of this issue.
8. J. O. McNally, G. H. Metson, E. A. Veazie and M. F. Holmes, Electron Tubes for the Transatlantic Cable System. See page 163 of this issue.

Repeater Production for the North Atlantic Link

By H. A. LAMB* and W. W. HEFFNER*

(Manuscript received September 20, 1956)

Production of submarine telephone cable repeaters, designed to have a minimum trouble-free life of twenty years, required many new and refined manufacturing procedures. Care in the selection and training of personnel, manufacturing environment, inspection, and testing, were of great importance in the successful attainment of the ultimate objective. Although quality of product has always been of major significance in Western Electric Company manufacture, building electronic equipment for use at the bottom of the ocean, where maintenance is impossible and replacement of apparatus extremely expensive, required unusual manufacturing methods.

MANUFACTURING OBJECTIVE

Late in 1952, the manufacture of flexible repeaters for the North Atlantic Link of the transatlantic submarine telephone cable system was allocated to the Kearny Works of Western Electric Company.

In accordance with established practice in initiating radically new products and processes, production of these repeaters was assigned to the Engineer of Manufacture Organization rather than to regular manufacture in the telephone apparatus shops. The job — to produce 122 thirty-six channel carrier repeaters and 19 equalizers capable of operating satisfactorily at pressures up to 6,800 pounds per square inch on the ocean floor, with minimum maintenance, for a period of at least twenty years. Initial delivery of repeaters was required for March, 1954, less than a year and a half after the project started.

GENERAL PHILOSOPHY

Quality has always been the prime consideration in producing apparatus and equipment for the Bell System. There is an economical breaking point, however, beyond which the return does not warrant the abnormal

* Western Electric Company.

expenditures required to approach theoretical perfection. The same philosophy applies to all manufactured commodities, be they automobiles, airplanes or telephone systems. In general, all of these products are physically available for preventive and corrective maintenance at nominal cost. With electronic repeaters at the bottom of the ocean, maintenance is impossible and replacement would be extremely expensive.

The general philosophy adopted at the inception of the project was to build integrity into the product to the limit of practicability. To do this, a number of fundamental premises were established, which form the foundation of all operations involved:

1. Manufacturing environment would be provided which, in addition to furnishing a desirable place to work, could be kept scrupulously clean and free from contamination.

2. The best available talent would be screened and selected for the particular work involved.

3. Wage payments would be based on day work, rather than on an incentive plan basis, because production schedules and the complexity of the operations did not permit the high degree of standardization essential to effective wage incentive operation.

4. A sense of individual responsibility would be inculcated in every person on the job.

5. Training programs would be established to thoroughly prepare supervisors, operators, and inspectors for their respective assignments before doing any work on the project.

6. Inspection, on a 100 per cent basis, would be established at every point in the process which could, conceivably, contribute to, or affect the integrity of the product.

PREPARATION FOR MANUFACTURE

Manufacturing Location

It appeared desirable to set up manufacture in a location apart from the general manufacturing area. Experience gained to date has satisfied us that this was the correct approach, since it provided a number of advantages:

1. Administration has been greatly facilitated by having all necessary levels of supervision located in the immediate vicinity of the work.

2. It was necessary for the people on the job to acquire and maintain a new philosophy of perfection in product, rather than a high output at an "acceptable quality level." This was easier at a separate location, since only one philosophy was followed throughout the plant.

3. Engineering, production control, service and maintenance organizations were located close to actual production and had no assignments other than the project.

4. The small plant, due to its semi-isolation, tends to produce a very closely knit organization and good teamwork.

A large number of manufacturing locations were examined and the one selected was a one-story modern structure in Hillside, New Jersey, which provided a gross area of 43,700 square feet.

The entire plant was air conditioned; in most cases, the temperature was controlled to minimum 73 degrees F, maximum 77 degrees F. The air was filtered through two mechanical and one electrostatic filters. Relative humidity was maintained at maximum 40 per cent in all but one area — the capacitor winding room — in which it was necessary to maintain maximum 20 per cent humidity to avoid mechanical difficulty with capacitor paper. While most of the air was recirculated, the air from the cafeteria, cleaning room, locker and toilet rooms was exhausted to the

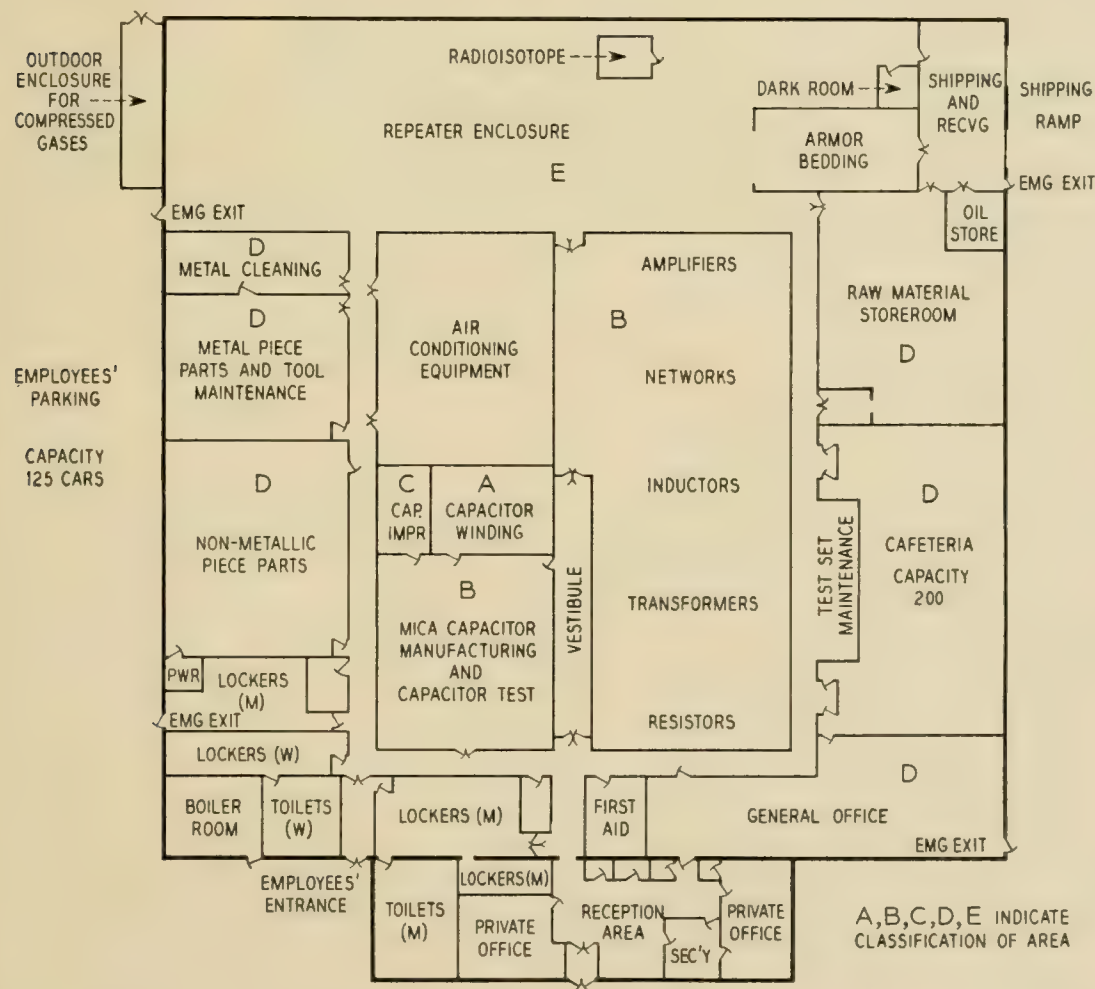


Fig. 1 — Plant layout.

outside atmosphere. Two separate air conditioning systems were in use. One, of 300 tons capacity, provided for most of the plant, while a smaller unit of 30 tons capacity served the capacitor winding, testing, and impregnating rooms. Each installation had its own air filtering and conditioning equipment.

Plant Layout

The plant layout is illustrated in Fig. 1. All working areas, with the exception of the repeater enclosure area, were individually enclosed, and walls from approximately four feet above the floor were almost entirely of reinforced glass. This arrangement facilitated supervision by other than first-line supervisors, who were located with the groups, and provided a means of viewing the operations by the many visitors at Hillside, without contaminating the critical areas or disturbing the operators.

Analysis of Design for Facilities and Operations

In analyzing the design for manufacture there were, of course, numerous instances where conventional methods and facilities were entirely adequate for the job. Since their inclusion would contribute little to this article, we shall confine the description to those cases which are new or unusual.

Collaboration with Bell Telephone Laboratories in Preparation of Manufacturing Information

Early in 1953 a coordination committee was established, consisting of representatives from the various Laboratories design groups and Western engineers, which met on a bi-weekly basis during the entire period preceding initial manufacturing operations. These meetings provided a clearing house for questions and policies of a general nature for this particular project and served to keep all concerned informed as to the progress of design and the preparations for manufacture.

It is customary, during the latter stages of development of any project at the Laboratories, for Western engineers to participate in the preparation of manufacturing information as an aid in pointing the design toward the most economical and satisfactory production methods and facilities. Since the decision to use the Bell System repeater in the Transatlantic system was based on the performance of the Key West-Havana installation, and the fact that changes in design would require further

trials over an extended period of time, only minor changes to facilitate manufacture were made. Further, since some experience had been gained by the Laboratories in producing repeaters for that installation, it was decided to "pool" effort in preparing the manufacturing process information, which is normally Western's responsibility. Close cooperation of the two groups, therefore, has resulted in the production of repeaters which are essentially replicas of those in the initial installation except for the internal changes necessary to increase transmission capacity from 24 to 36 channels.

Other Western Electric Locations and Outside Suppliers

During the development work on the Key West-Havana repeaters, the Hawthorne Works of Western Electric had furnished the molybdenum-permalloy cores for certain inductors, the Tonawanda Plant had furnished mandrelated resistance wire, and the Allentown Plant had fabricated the glass seal subassemblies. Since the experience gained in this development work was extremely valuable in producing the additional material required for the Transatlantic system and since the facilities for doing the work were largely available, these various locations were asked to furnish similar material for the project. Although the Kearny Crystal Shop had not been involved in the Key West-Havana project, arrangements were made there to make the crystals for this project, since facilities were available, along with considerable experience in producing precision units.

Subcontracted Operations

While it was believed, initially, that all component parts for repeaters should be manufactured by Western Electric, critical analysis indicated that it was neither desirable nor economical in certain cases. One of the outstanding examples in this category is the hardened and ground chrome-molybdenum steel rings that constitute the strength members in the repeater and sustain the pressures developed on the ocean bottom. Purchasing the many large and varied machine tools and associated heat treating equipment necessary to produce these parts would have required a substantial capital expenditure and additional manufacturing space. Arrangements, therefore, were made with a highly qualified and well equipped supplier to produce the rings, using material furnished by Western, which had been previously inspected and tested to very stringent requirements.

The situation attending the manufacture of a relatively small number of comparatively large copper parts used in the rubber and core tube seals was much the same. Here, again, the large size machine tools and additional manufacturing space, required for only a short time, would have increased the over-all cost of the project considerably. These parts, therefore, were subcontracted in the local area and inspection was performed by Hillside inspectors.

A safeguard, in so far as integrity is concerned, was provided by the fact that these were individual parts that could be reinspected at the time of delivery. No subassembly operations that might possibly result in oversight of a defect, were subcontracted.

Manufacturing Conditions

Two major problems confronted us in planning the manufacture of repeaters. First, to produce units that were essentially perfect; and second, to prevent the contamination of the product by any substance that might degrade its performance over a long period of time. In approaching both of these objectives, it was realized that the product had a definite economic value which the cost of production should not exceed. In many cases, therefore, it was necessary to rely on judgment, backed by considerable manufacturing experience, in determining when the "point of no return" had been reached in refining processes and practices.

The initial approach to this phase of the job was to classify, with the collaboration of Bell Telephone Laboratories, all of the manufacturing operations involved as to the degree of cleanliness required. In setting up these criteria, it was necessary to evaluate the importance of contamination in each area and the practicability of eliminating it at the source or to insure that whatever foreign material accumulated on the product was removed.

A representative case is the machining of piece parts. While the shop area is cleaner, perhaps, than any similar area in industry, the very nature of the work is such that immediate contamination cannot be avoided since material is being removed in the form of chips and turnings, and a water soluble oil is used as a coolant. In this instance, however, the parts can be thoroughly cleaned and their condition observed before leaving the area. Conversely, in the case of an operation such as the assembly of paper capacitors into a container which is then hermetically sealed, it is vitally necessary to insure that both the manufacturing

area and the processes are free from, and not conducive to producing, particles of material which are capable of causing trouble.

The various classifications established for the production areas include specific requirements as to temperature, relative humidity, static pressure with respect to adjacent areas, cleanliness in terms of restrictions on smoking and the use of cosmetics and food, and the type and use of special clothing.

Special Clothing

Employees' clothing was considered one of the most important sources of contamination for two reasons; first, for the foreign material that could be collected upon it and carried into the manufacturing areas, and second, that various types of textiles in popular use are subject to considerable raveling and fraying.

After considerable study of many types of clothing for use in critical areas, the material adopted was closely woven Orlon, which has proved to be acceptably lint-free. The complete uniform — supplied at no cost to employees — consists of slacks and shirts for both male and female employees, Orlon surgeon's caps for the men and nylon-visored caps for the women. In addition shoes, without toecap seams, were provided. Nylon smocks were furnished to protect the uniforms while employees moved from locker rooms to the entrance vestibule. Two changes of clothing were provided each week, and the laundering was done by an outside concern.

Employees to whom this special clothing was issued were paired for locker use. Both kept their uniforms and special shoes in one locker and their own clothes and shoes in the other. This prevented the transfer to the uniforms of any foreign material that might exist on the street clothing. At the entrance vestibule to the A, B, and C areas (Fig. 1) the employees were required to clean their shoes in the specially designed facilities provided and to wash their hands in the wash basins installed for this purpose. Smocks were then removed and hung on numbered hooks that line the walls at the end of the vestibule. Employees were then permitted to go to their work positions within the inner areas. At any time that it was necessary for employees to leave the work areas for any purpose, they were required to put on their smocks in the vestibule and upon their return, to go through the cleaning procedure again.

Employees in the other areas were provided only with smocks, mainly for the protection of their clothes since the work involved could soil or stain them but could not be contaminated from the clothing.

Cleaning

Schedules were established for cleaning the areas at regular intervals, the frequency and methods depending upon the type of manufacturing operations and the activity. Usually, the vinyl plastic floors were machine scrubbed and vacuum dried. Walls, windows and ceilings were cleaned by hand with lint-free cloths. Manufacturing facilities such as bench tops, which were linoleum covered, were washed daily. Test sets, cabinets, test chambers and bench fixtures were also cleaned daily. Hand tools were cleaned at least once a week by scrubbing with a solution of green soap, rinsing in distilled water, followed by alcohol and then dried in an oven.

Dust Count

Since it was impossible to determine what contaminating material in the form of air-borne particles might be encountered from day to day, and what the effect might be during the life of the repeaters, the general approach to this problem was to control, so far as possible, the amount of dust within the plant.

In order to verify, continuously, the over-all effectiveness of the various preventive measures, dust counts were made in each classified area at daily intervals, using a Bausch and Lomb Dust Counter. This device combines, in one instrument, air-sampling means and a particle-counting microscope. Over a two-year period it has been possible to maintain, in certain areas, a maximum dust count of between 2,000 and 3,500 particles per cubic foot of air with a maximum size of 10 microns. Control checks, taken outside the building at the employees' entrance, generally run upwards of 25,000 particles per cubic foot, a good portion of which are of comparatively large size.

PRODUCTION AND PERSONNEL

Equipping the plant, obtaining and installing facilities, and selecting and training personnel proceeded on a closely overlapped basis with receipt and analysis of Bell Telephone Laboratories' product design information. Because of the critical nature of the product, provisions were made not only for the most reliable commercially available utilities and services, but also for emergency lighting service in some areas. Maintenance and service staffs had to be built up rapidly as the supervisory and manufacturing forces were being developed.

"Qualification" of All Personnel

Before employees were assigned to production work they were required to pass a qualification test established by the inspection organization to demonstrate satisfactory performance. Programs were, therefore, set up for "vestibule" training and qualification of new employees. This activity was carried on by full-time instructors who had been trained by Western and Bell Laboratories engineers. Training was carried out in two stages:

1. (a) The employee received instruction and became acquainted with equipment and requirements. (b) A practice period in which the employee developed techniques and worked under actual operating conditions, with all work submitted to regular inspection.

2. A qualification period in which the employee was required to demonstrate that work satisfactory for project use could be produced.

The main objective during stage 1 was progressive quality improvement and in stage 2 the maintenance of a satisfactory quality level over an extended period of time. Employees made a definite number of units at acceptable quality levels in order to qualify. The number of units required for training varied with the type of work and the ease with which it was mastered.

All personnel were required to pass qualification tests before being assigned to production work and were restricted to that work unless trained and qualified for other work. Employees trained on more than one job were requalified before being returned to a previous assignment.

Records of the performance of individual operators started in the training stage were continued after the employees were assigned to production work. The performance record of the operators was based on results obtained during the inspection of their work, while that of the inspectors was based on special quality accuracy checks of their work.

Personnel Selection

It was apparent that the new manufacturing techniques, including the cleanliness and quality demands, would necessitate that all shop supervisors and employees be very carefully selected. It also appeared (and this was subsequently confirmed) that after the careful selection and training of supervisors, long training periods would be required for specially selected shop employees.

In selecting first line shop supervisors, such factors as adaptability, personality, and ability to work closely with the engineers were of paramount importance. For the parts and apparatus included in their re-

sponsibility, they were required to thoroughly learn the design, the operations to be performed, the facilities to be used, the data to be recorded, the cleanliness practices to be observed — and in most cases, prepare themselves to be able to do practically all of the operations, because subsequently they had to train selected operators to perform critical operations to very high quality standards under rigidly controlled manufacturing conditions. As shop supervisors and employees were assigned to the manufacture of repeaters, they were thoroughly indoctrinated in the design intent and the new philosophy of manufacture.

Standard ability and adaptability tests were used in a large number of cases to assist in proper selection and placement of technicians. Tests for finger and hand dexterity; sustained attention; eyes, including perception and observation; and reaction time of the right foot after a visual stimulus. (The latter test was relatively important for induction brazing operations.) Other requisite considerations were a high degree of dependability and integrity, involving intellectual honesty and conscientious convictions; capability of performing tedious, frustrating, and exasperating operations against ultra-high quality standards, verifying their own work; perseverance and capability to easily adapt to changes in assignment and occupation or the introduction of design changes. We considered whether or not they would stand up under “fishbowl” operations, wherein they would receive a considerable amount of observation from high levels of Western Electric Company and Bell System management and other visitors. Also, could they duplicate high quality frequently after qualifying for a particular operation?

During the period of repeater manufacture, the number of employees rose from less than 50 in January, 1954, to a maximum of 304 by February, 1955, after which there was a gradual reduction to a level of about 265 employees for six months and then a gradual falling off as we were completing the last of the project. In the period from May to December, 1954, between 30 and 45 employees were constantly in training prior to being placed on productive work. During 1955 this decreased to practically no employees in training during the midpart of the year and thereafter training was required merely to compensate for a small labor turnover and employee reassignment. It is significant that labor turnover was very low and attendance was exceptionally good during the life of the Hillside operations.

Personnel Training

The original plan, which was generally followed, was to prove in the tools for each phase of the job, followed by an intensive program of train-

ing. Indoctrination of laboratory technicians could be considered as "vestibule training" in that they were acclimated to the area and conditions, given oral instruction in the work, then given practice materials and demonstrations and, when qualified, were started on making project material. To do this, extra supervisors were required at the beginning of the job. A supervisor trained a few employees, qualified some of them, and began work on the project. Another supervisor was then required to train additional employees who, as they became qualified, were transferred to the supervisor responsible for making project apparatus. Additional testing of the employees, instruction and reinstruction and, in some cases, retraining were required. In practically all cases, we were able to fit an employee selected for work at Hillside into some particular group of operations. The extra emphasis on selection and training created a well-balanced team that later resulted in considerable flexibility. During all of this training our supervisors worked closely with engineers and inspectors who understood the design intent and the degree of perfection required.

At the beginning, each technician was trained for only one operation of a particular job, such as (1) winding Type X capacitors or (2) impregnating all paper capacitors or (3) winding Type Y transformers and so became an expert on this one operation. Later, the tours of duty for many technicians were broadened to cover several operations.

Communications

To keep employees informed, we occasionally assembled the entire group, presenting informative talks on current production plans and our future business prospects. Motion pictures were shown of the cable laying ships and the operations of cable splicing and cable laying. A display board, showing all of the repeater components, was mounted on the wall of the cafeteria. This informed the operators just where the parts were used in apparatus; also, just where their products went into the wired repeater unit, and how all electrical apparatus was enclosed against sea pressure in the final repeater. In small groups, all of the employees at Hillside were given a short guided tour of the plant to see the facilities and hear a description of the operations being performed in each area. These communications were extended to everyone at the Hillside Plant, including those who did not work directly on the product. It was our conviction that the maintenance men, boiler operators, oilers, station wagon chauffeur, janitors, and clerical workers in the office were all interested and could do a better job if kept informed of the needs and progress of the project.

Scheduling

Capacity was provided at the Hillside Shop to manufacture a maximum of 14 repeaters in a calendar month. This envisioned 6-day operation with some second and third shift operations; due allowance was made for holidays and vacations, so that the annual rate would be approximately 160 enclosures per year. (An enclosure is either a repeater or an equalizer.)

Some of the facilities and raw materials were ordered late in 1953. This ordering expanded early in 1954 and continued through 1955 to include parts to be made by outside suppliers and the parts and apparatus to be made at Hillside. Apparatus designs were not all available at the beginning of the job, and the ultimate quantities required were also subject to sharp change as the project shaped up, thus further complicating the scheduling problem.

Because of the time and economic factors involved, coupled with the developmental nature of the product and processes, one of the most difficult and continuing problems was the balancing of production to meet schedules. For this task, we devised "tree charts" for the apparatus codes and time intervals in each type of repeater or equalizer for each project. Each chart was established from estimates of the time required to accomplish the specified operations and the percentage of good product each major group of operations was expected to produce.

RAW MATERIALS

Many of the specifications were written around the specific needs of the job and embodied requirements that were considerably more stringent than those imposed on similar materials for commercial use. As a result, it was necessary for many suppliers to refine their processes, and, in some cases, to produce the material on a laboratory basis.

One example is the container, or repeater enclosure, which consists, in part, of a seamless copper tube approximately $1\frac{3}{4}$ inches in diameter having a $\frac{1}{32}$ -inch wall and approximately 8 feet long. This material was purchased in standard lengths of 10 feet. The basic material was required to be phosphorous deoxidized copper of 99.80 per cent purity. The tubing, as delivered, had to be smooth, bright, and free from dirt, grease, oxides (or other inclusions including copper chips), scale, voids, laps, and slivers. Dents, pits, scratches, and other mechanical defects could not be greater than 0.003 inch in depth. The tubing had to be concentric within 0.002 inch and the curvature in a 10-foot length not exceed $\frac{1}{2}$ inch to facilitate assembly over the steel rings.

Only one supplier was willing to accept orders for the tubes, and only on the basis of meeting the mechanical requirements on the outside surface. To establish a source of supply, it was necessary to accept the supplier's proposal on the basis that some of the tubes produced could be expected to meet requirements on the inside as well as the outside surface. Inspection of the inside surface was performed with a 10-foot Bore-scope.

The supplier then set aside, overhauled, and cleaned a complete group of drawing facilities for this project. In addition, a number of refinements were made in lubrication and systematic maintenance of tools. After all refinements were made and precautions taken, however, the yield of good tubes in the first 400 produced was less than 1.0 per cent. Consultations with Western and Bell Laboratories' engineers, and with the supplier's cooperation, raised the yield to approximately 50 per cent.

Procurement of satisfactory mica laminations for capacitors introduced an unusual problem. The best grade of mica available in the world market was purchased which the supplier, under special plant conditions, split and processed into laminations. Despite care in selection and processing, only 50 per cent of the 250,000 laminations purchased met the extremely rigid requirements for microscopic inclusions and delaminations, and less than 8 per cent survived the capacitor manufacturing processes.

A large number of the parts, and the most complex, are made from methyl-methacrylate (Plexiglass). At the time manufacture began, there was little, if any, experience or information available on machining this material to the required close tolerances and surface finish. Consequently, considerable pioneering effort was expended in this field before satisfactory results were obtained.

The methacrylate parts cover a wide range of size and complexity — from $1\frac{1}{2}$ -inch diameter by $4\frac{7}{8}$ -inch long tubular housing to tiny spools $\frac{1}{8}$ -inch diameter and $\frac{1}{16}$ -inch long. Most of the parts are cylindrical in shape with some semicylindrical sections that must mate with other sections to form complete cylinders. Others have thin fins, walls, flanges and projections. Five representative parts are shown in Fig. 2.

Methyl-methacrylate has a tendency to chip if tools are not kept sharp and care is not used in entry or exit of the tool in the work, particularly in milling. In some cases, it is necessary, with end-milling, to work the cutter around the periphery of the area for a slight depth so that subsequent cuts will not break out at an unsupported area. Normally, with a sharp cutter and a 0.010-inch finish cut, and a slow feed, chipping will not result. High-speed steel tools with zero rake were used for turning

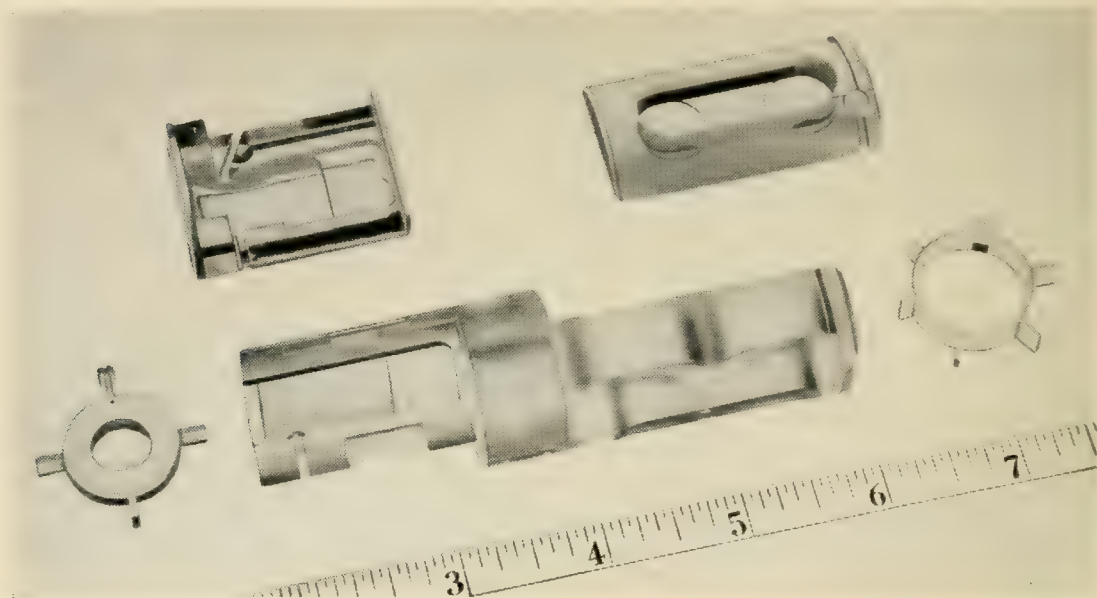


Fig. 2 — Methyl methacrylate parts.

and boring operations. Standard high-speed milling cutters and end mills were used for milling except for the cutting edges, which are honed to a fine finish. A clearance angle of 7 degrees for milling and 10 degrees to 15 degrees for lathe work was found most satisfactory. In lathe work, the general rule was light feeds (0.003 inch–0.005 inch) and small depth of cut. However, the depth of cut could be safely varied over a wide range depending upon many factors, such as type of part, quality of finish, machine and tool rigidity, effective application of coolant, and tooling to support and clamp the part. In one operation of boring a $1\frac{3}{16}$ -inch diameter by $4\frac{1}{8}$ -inch deep blind hole within ± 0.002 inch, the boring terminates in simultaneously facing the bottom of the hole square with its axis. A cut $\frac{1}{32}$ -inch deep with a light feed was taken with a specially designed boring tool with the coolant being fed through the shank to the cutting edge. All completely machined parts were annealed for 12 hours at 175° F.

HIGHLIGHTS IN ASSEMBLY AND BRAZING

Repeater units are encased in hardened steel rings which previously had been tested at 10,000 pounds per square inch hydraulic pressure. This pressure is approximately 50 per cent higher than the greatest pressure expected at ocean bottom. The steel rings were encased in a copper sheath and closed at each end with a glass-to-Kovar seal, with the central conductor coming through the glass to the outside. The copper sheath was then shrunk to the steel rings and glass seals using 6000 pounds per

square inch hydraulic pressure, and the glass seal was then high-frequency brazed to the copper sheath.

To keep the ocean bottom pressure off the glass seals and also to terminate the cable insulation, a rubber seal is brazed in to the copper container tube adjacent to each glass seal. This rubber seal consists of rubber bonded to brass, which has been brazed to the copper portion of the seal. The rubber terminates in polyethylene through five steps of compounds containing successively less rubber and more polyethylene. The polyethylene can be readily bonded by molding to the polyethylene insulation of the cable. The central conductor passes through a central brass tube in the rubber seal, which is also bonded to the rubber.

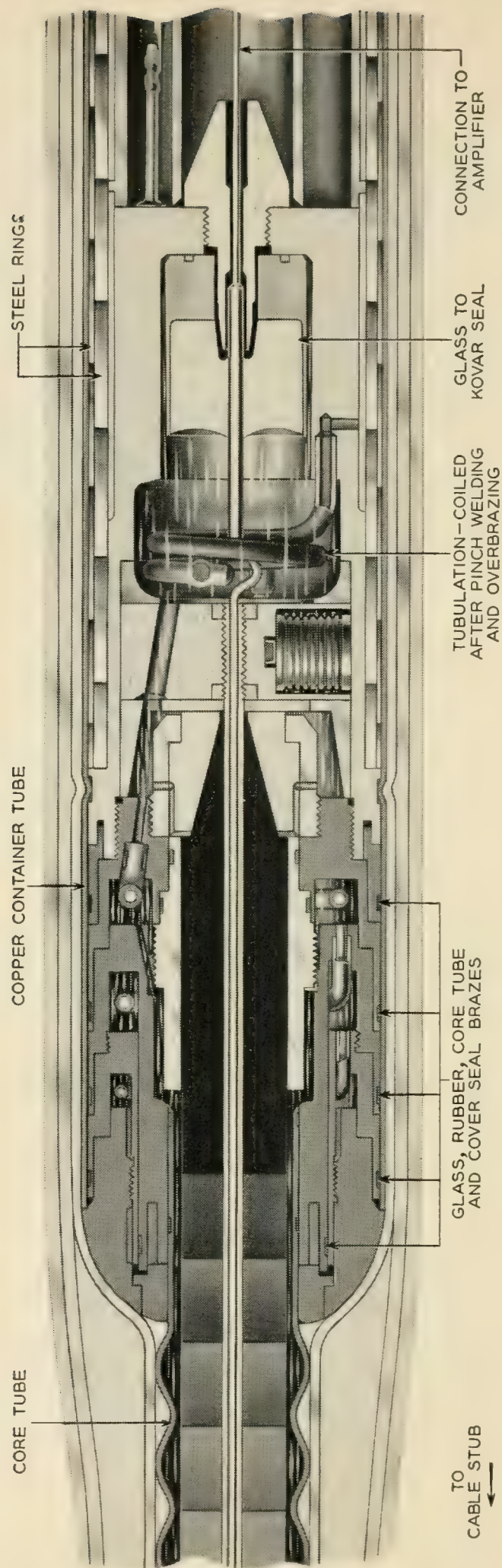
To protect the rubber seals from the deleterious effects of salt water immersion for long periods of time, a copper core tube is brazed over each rubber seal. The core tube is arranged to equalize the pressure inside and out when submerged at ocean bottom pressure. This is accomplished with a bulge of neoprene filled with polyisobutylene, on the far end of the core tube, which transmits the pressure to the inside of the core tube seal.

To make doubly sure that no salt water reaches the rubber seal, a copper cover is brazed into the container outside the core tube connector on each end. This cover is also brazed to the core tube connector. The interstice between each of the above four seals is filled with polyisobutylene, which is viscous and inert and has very good insulating qualities.

Each end of the repeater closure (Fig. 3) contains five successive brazed joints. Any one of these ten brazes, if not perfect, could cause the loss of the repeater closure and jeopardize the entire repeater. All of these brazes were made with the repeater in a vertical position to insure an even distribution of the brazing alloy fillet around the joint.

An upending device was provided at the pit brazing location to raise the repeater on its carrier to a vertical position with either end up and move it into position for brazing. The repeaters were brought into the brazing area on an overhead monorail and an electric hoist. The shorter repeater assemblies, before core tube and cable stub assembly, were upended by hand and brazed from a raised platform.

It was necessary to make all of these brazes by high-frequency induction heating, since the heat must be intense, contained within a very narrow band, evenly distributed, and the area protected from oxidation by a somewhat reducing atmosphere. The heat must be very intense since the time interval for the shortest braze was 10 seconds maximum and the longest was 30 seconds. A large part of the heat was dissipated



F. 3 --- Repeater seals.

by being conducted at a high rate from the copper parts to the water in the cooling jackets used to contain the heat in a very narrow band.

Circulating cooling water within a jacket prevented heat from being conducted down the copper container tube to the preceding seals or to the repeater unit. This water-cooled jacket was positioned only $\frac{3}{8}$ inch below the inductor, and the water was in intimate contact with the container tube, which is sealed off at both ends with rubber "O" rings. In addition, for the glass seal braze, the glass inside the seal cavity was kept covered with water during the heat cycle. The water was fed in and siphoned out to a constant level which was kept under observation by the operator and the inspector to make sure that the glass was covered at all times. The rubber seal was also water jacketed on the inside of the seal to prevent deleterious effects of the heat on the rubber insulation around the central conductor. The inner cover braze was quenched before the 10-second maximum interval had expired to insure that the heat did not penetrate to the polyisobutylene at a sufficient rate to deteriorate it or the rubber inside.

Distribution of the heat around the container tube at the braze area was controlled by locating the work in the inductor so that the color came up essentially evenly all the way around and at the proper level to bring a fillet up to the top of the braze joint within the allowable time limit. The time limit was determined by experiment so that none of the previously assembled parts were damaged by the heat. This determination of the proper heat pattern and the prevention of overheating required the development of considerable skill on the part of the operator. The variables encountered made it essential to rely on an operator to control the heat rather than to utilize the timer with which the induction heating equipment is normally controlled.

The area to be heated for brazing was protected from oxidation by enclosing it in a separable transparent plastic box and flooding the interior with a gas consisting of 15 per cent hydrogen and 85 per cent nitrogen. This atmosphere is somewhat reducing and not explosive. The brazing surfaces of the parts were chemically cleaned immediately before assembly and extreme care was exercised to keep them clean until brazed.

The container tube was shrunk to the respective glass, rubber, core tube, and cover seals using hydraulic pressure so that the surfaces to be brazed and the brazing alloy were in intimate contact within the brazing area. If the parts were clean and kept from oxidizing by the protective atmosphere, the alloy would flow upward by capillary action and form a fillet around the top of the seal, impervious to any leak.

The braze in each case was then leak tested with a helium mass-spec-

trometer type leak detector. A gas pressure of helium at least 25 per cent greater than the maximum pressure to be encountered at ocean bottom was used. In addition, a radioisotope was used to test the effectiveness of the final tubulation pinch welds and overbrazes which were kept open for the leak tests under high pressure helium. These tests were made with water pressure about 25 per cent greater than the maximum ocean bottom pressure.

The completed repeater was inserted in a chamber 80 feet long; the chamber was then filled with water and the pressure raised to 7,500 pounds per square inch and held at that pressure for at least 15 hours. At the end of this period the closure had to show no sign of crushing or leaking.

The repeater unit sealed in the closure must be extremely dry to function properly. Any water vapor which might remain after the closure is sealed, or enter during the estimated 20-year minimum life, must be scavenged. A sealed desiccator with a thin diaphragm was, therefore, assembled into the repeater unit sections. After completely drying and sealing the repeater unit except for one tubulation, the diaphragm of the desiccator was ruptured by dry nitrogen pressure and with the enclosure filled with dry nitrogen the final tubulation was immediately sealed off. To insure that the diaphragm was actually broken, a microphone was strapped to the outside of the repeater over the location of the desiccator and a second microphone arranged at the end of the closure to pick up background noises. A pen recorder was used to record the sound from the two microphones and also the change in nitrogen pressure. Three simultaneous pips on the chart gave definite indication that the diaphragm had ruptured and that the desiccant had been exposed to the internal atmosphere of the repeater.

QUARTZ CRYSTAL UNITS MANUFACTURED AT KEARNY

The primary purpose of the crystal unit is to provide the means of identifying and measuring the gain of each repeater in the cable. This basic crystal design is in common usage. The exacting specifications for this application, however, imposed many problems and deviations from normal crystal manufacturing processes.

Raw Quartz was specially selected for this crystal unit. The manufacturing process of reducing the quartz to the final plate followed the recognized methods through the roughing operations. Due to the rigid end requirements, the finishing operations were performed under laboratory conditions. Angular tolerances were one-third of normal limits. No evidence of surface scratches, chipped edges or other surface imper-

fections visible under 30X magnification were permitted. This resulted in a process shrinkage five times that experienced in normal crystal plate manufacture.

In this use, the crystal units were required to meet performance tests at currents as low as one-thousandth of a microampere — far below the current values usually encountered. Improved soldering techniques had to be developed for soldering the gold plated phosphor bronze and nickel wires used, because it was found that the electrical performance of the units was directly related to the quality of soldered connections.

Although one-seventh of Western's production of quartz crystal units are in glass enclosures, the applicable techniques in glass working required a complete revision. Glass components such as the stem and bulb purchased from established sources were found to be far below the standard required for this crystal unit. For example, the supplier of the glass tubing used in the manufacture of stems was required to meet raw material specifications that embodied coefficient of thermal expansion, softening point of glass, density, refractive index, and volume resistivity. The glass stems made from this tubing by regular manufacturers were found unacceptable and the processes used by these sources could not be readily adapted to meet the desired specifications. The glass stems contained four lead wires made from 30-mil Grade "A" nickel wire butt welded to 16-mil light borated Dumet wire. To assure the quality of the metal to glass seal, each wire was inspected under 30X magnification for tool marks and other surface imperfections. The finished stem assemblies were inspected under 30X magnification for dimensions, workmanship, cleanliness and minute glass imperfections, then individually stored in a sealed plastic envelope.

The glass bulb in this crystal unit is known as the T921 design commonly used in the electron tube industry. The high quality required, however, made 100 per cent inspection necessary. Examination under 30X magnification resulted in rejected bulbs for presence of scratches, open bubbles, chips and stones. Physical limits for inside and outside diameters as well as wall thickness were causes for additional rejects. Only one per cent of the commercial bulbs were found acceptable, and these were also stored in a sealed plastic envelope.

The final major assembly operation consisted of sealing the glass bulb to the stem which had had the crystal sub-assembly welded to the nickel wires. The techniques for "sealing in" used in quartz crystal or electron tube manufacture were unsuited. Two important factors in this crystal unit, which required the development of new processes, were the proximity of soft soldered connections to the sealing fires and the demands

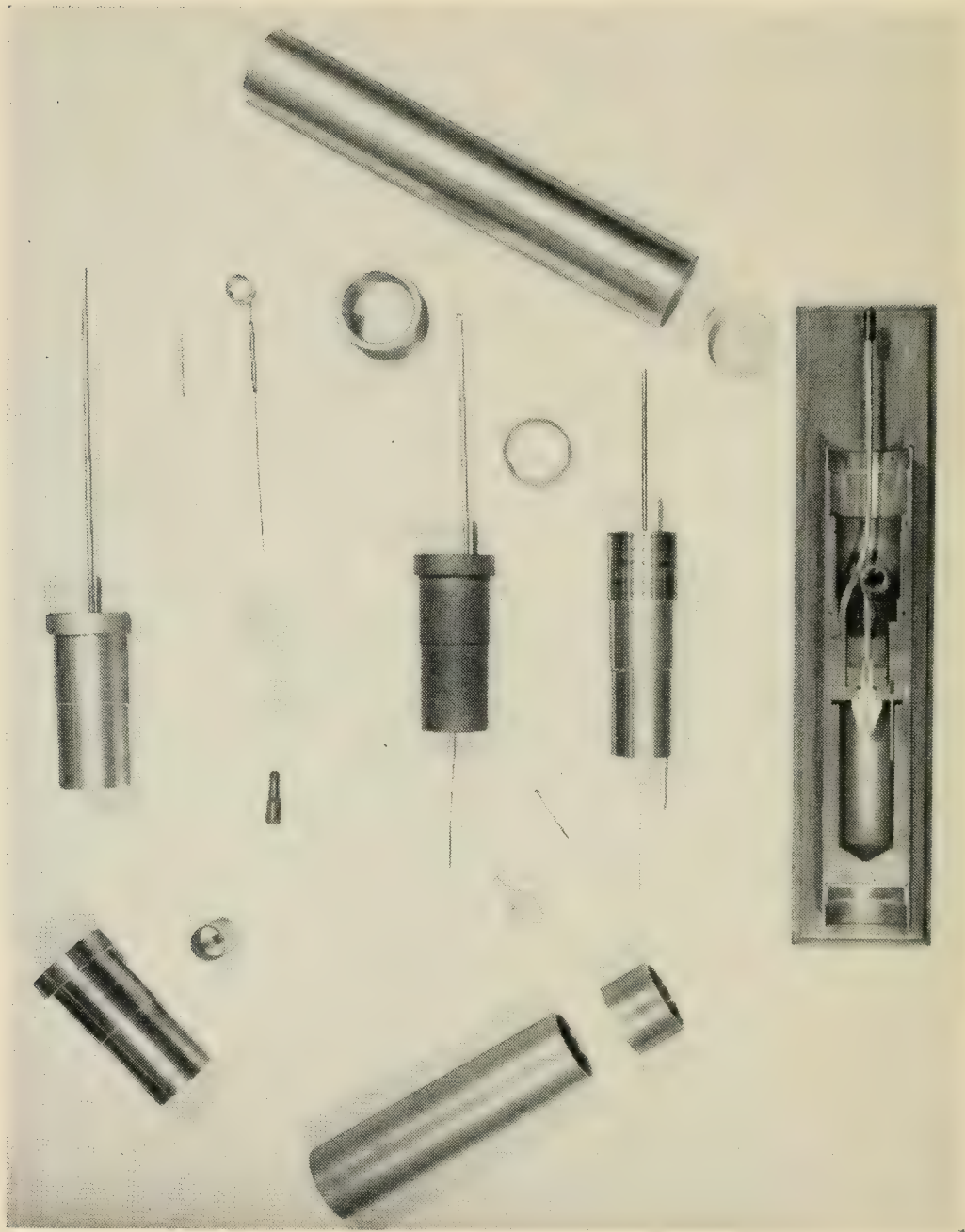


Fig. 4 — Glass to Kovar seal.

that the glass seal contain a minimum of residual tensile stress. These two problems were resolved collectively by performing the sealing operation on a single spindle glass sealing machine. Accurate positioning of the glassware and sealing fires, together with precise timing and temperature controls, achieved the desired results.

Evaluation of residual stresses were made by inspections using a polarimeter and by a thermal shock test. The maximum safe stress was established at 1.74 KG/mm^2 . The thermal shock test required successive immersion of the unit in boiling water and ice water. The electrical characteristics of these units exceeded all others made previously by Western Electric. The ratio of reactance to effective resistance ("Q") was greater than 175,000 — twice that ever previously produced and 17 times that required in the average filter crystal.

Stability for frequency and resistance was assured by a 28-day aging test. During this period, precise daily resonant frequency and resistance measurements were recorded against temperature within 0.1°C . The maximum permissible change was 0.0005 per cent in frequency and $+5$ per cent to -10 per cent in resistance.

GLASS SEALS MANUFACTURED AT ALLENTOWN

The glass seal used to close each end of the container for the repeaters and equalizers is manufactured at the Allentown Works of the Western Electric Company.

The unit is essentially a glass bead-type seal. It insulates the central conductor of the repeater from the container and serves as a final vapor barrier between the cable and the interior of the repeater. As such, it backs up several other rubber and plastic barriers as shown in Fig. 3.

Fig. 4 shows the various components, subassemblies, and a cross-section of the unit. The unit consists of the basic seal brazed in the Kovar outer shell, to which is brazed a copper extension provided with two brazing-ring grooves. One of these grooves is used in brazing the seal, along with support members, into a length of container tubing in the same manner as the seal is ultimately brazed into the repeater. Packaging of the seal in this manner was necessary to pressure test the seal. Under test, in a specially constructed chamber 10,000 psi of helium gas pressure was applied to the external areas of the packaged glass seal and a mass spectrometer type leak detector was connected through the tubulation to the internal cavity of the packaged unit. In this manner, the interface of the glass to metal seal, the brazed joints, and the porosity of the metal were checked for leakage. The unit is left in this package for delivery to provide protection during shipment. Before the seal could be used,

it was machined from the package by cutting the copper extension to length, leaving the second groove for use in brazing the seal to the repeater and removing the container tubing and the support members.

The basic seal consists of the cup, central conductor and glass. The cup (smaller cylindrical item in the upper lefthand corner of Fig. 4) was machined from Kovar rod. The wall of the cup is tapered from a thickness of 0.025 inch at the base to 0.002 inch at the lip. The last 0.006 inch of the lip is further tapered from this 0.002 inch to a razor edge. The internal surface is better than a 63-micro-inch turned finish and was also liquid honed to give it a uniform matte finish. The central conductor (slim piece in the upper right-hand corner of Fig. 4) was also machined from Kovar rod. Both the cup and central conductor were further processed by pickling, hypersonically cleaning in deionized water, and decarburizing. The glass, a borosilicate type of optical quality, was cut from heavy walled tubing. The glass tubing was hand polished, lapped and etched to remove surface scratches, and to arrive at the specified weight. It was also fire polished and hypersonically cleaned to remove all traces of surface imperfections and to assure maximum cleanliness.

In order to make the basic glass seal, the metal parts had to be oxidized under precisely controlled conditions. For the oxidizing operation, a suitable fixture was loaded with brazed shell-cup assemblies, central conductor assemblies, and a Kovar disc, which had been prepared in precisely the same manner as the cups and central conductors. The disc was carefully weighed before and after oxidizing and the increase in weight divided by the area involved yields the weight gain due to oxidation for each run. Limits of 1.5 to 2.5 milligrams per square inch of oxide were set. This operation was performed by placing the loaded, sealed retort, through which passed a metered flow of dried air, into a furnace for a specified time-temperature cycle.

In the glassing operation the oxidized shell assembly, the carbon mold and the central conductor were placed in a fixture and held in the proper relationship. The carbon mold served to support the glass, while it was being melted, in that section between the cup and central conductor where the glass was normally unsupported. The prepared cut glass tubing was loaded into the Kovar cup and the fixture was sealed into the retort. During the glassing cycle, a constant flow of nitrogen passed through the retort to provide an atmosphere which minimized any reduction or further oxidation of the already carefully oxidized parts. After the proper purging period, the retort was placed in the furnace. In the furnace, the glass melted and formed a bond with the oxidized Kovar of the cup and

central conductor to form the seal. After the specified temperature-time cycle, the retort was removed from the furnace, allowed to partially cool and then placed into an annealing oven.

Vertical furnaces and retorts were used for brazing, decarburizing, oxidizing and glassing. By varying the type of gases flowing into the retorts, atmospheres which are reducing, oxidizing, or neutral were obtained. To provide maximum uniformity of process, separate retorts and holding fixtures were provided for operations involving hydrogen and for air-nitrogen operations, so that a retort or a fixture used for hydrogen treatments was never used for oxidizing or glassing.

PILOT AND REGULAR PRODUCTION

We called our first efforts *Practice Parts and Training*; the next we called Pilot Production. Next, certain items identified as *Trial Laying Repeaters and Oscillators* were manufactured for use in "proving in" the ship laying gear. To prove in manufacturing facilities, a few unequipped housings were made without the usual electrical components normally in a repeater. Similarly, each of the apparatus components and parts required exploratory and pilot effort before regular production could be undertaken.

As might be expected, the manufacturing yield of components meeting all requirements was very low during the early stages of the undertaking. However, substantial improvement was brought about as experience was gained. Comments on some of the production problems, highlights, and yield results, follow.

Paper Capacitors were manufactured only after painstaking qualifying trials and tests had been performed on each individual roll of paper. Cycling and life testing, procurement of acceptable ceramic parts and gold-plated tape and cans, selection and matching of rolls of paper for winding characteristics, and similar problems, all had to be completely resolved to a point of refinement previously unattempted for telephone apparatus.

Composite percentage yield for all operations on paper capacitors is shown in Fig. 5. Yield is shown as the ratio of finished units of acceptable quality to the number of units started in manufacture.

Mica Capacitors were made from only the most meticulously selected laminations, as mentioned earlier. Even the best mica is particularly susceptible to damage in processing. In spite of experience and knowledge of this, the multiple handling of the laminations contributed an unusually high material shrinkage as each separate lamination needed to be cleaned, then handled individually many times through the proc-

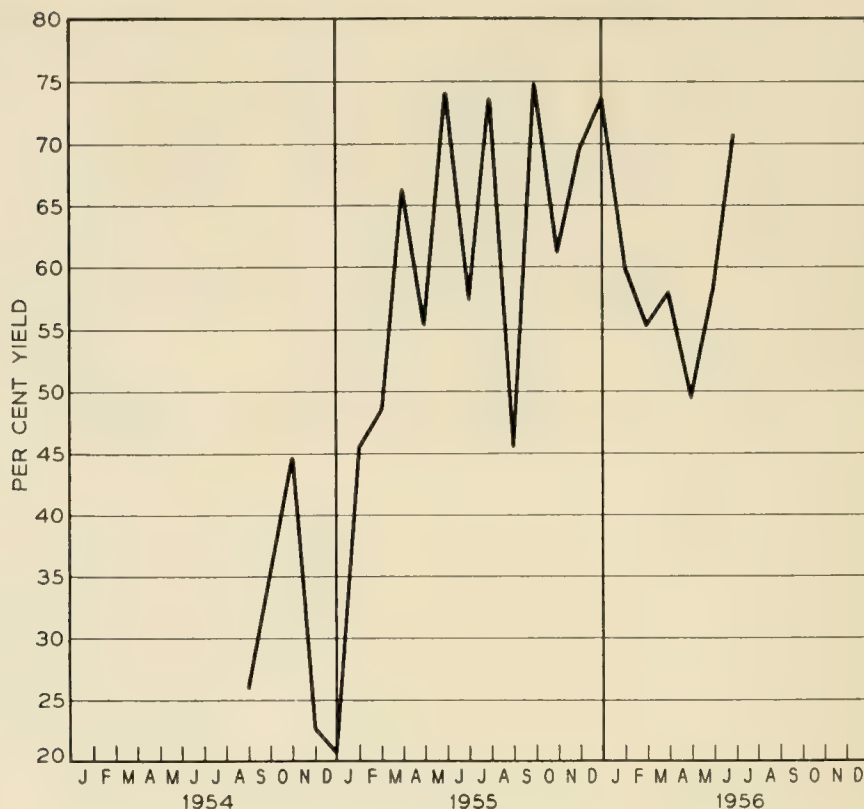


Fig. 5 — Paper capacitor yield.

esses. The art of silk screening was applied to deposit silver paste in a specific area or areas on each side of a lamination. A sharply defined rectangular area was required so that when superimposed one over another the desired capacitance would be obtained. Cementing of mica laminations onto machined methacrylate forms presented some additional problems through the bowing of the mica laminations as the cement cured. Obtaining screens that would give the proper length and width dimensions for the coated area, was another problem. A silk screen woven of strands of silk obviously limits, by the diameter of the threads, the extent to which the dimensions of an opening may be increased or decreased. Beryllium copper U-shaped terminals were used to clamp the layers of mica together into a stack. Control of the pressure used in crimping these terminals was found to be very critical in view of the exceptionally tight limits on capacitance and stability. Fig. 6 shows the composite yield at various times for all mica capacitors.

Resistors. There were three designs of ceramic resistors, which were resistance-wire wound on ceramic spools. These were intended to be assembled into the hole inside the core tube on which the paper capacitors were wound. Special winding machines equipped with binocular

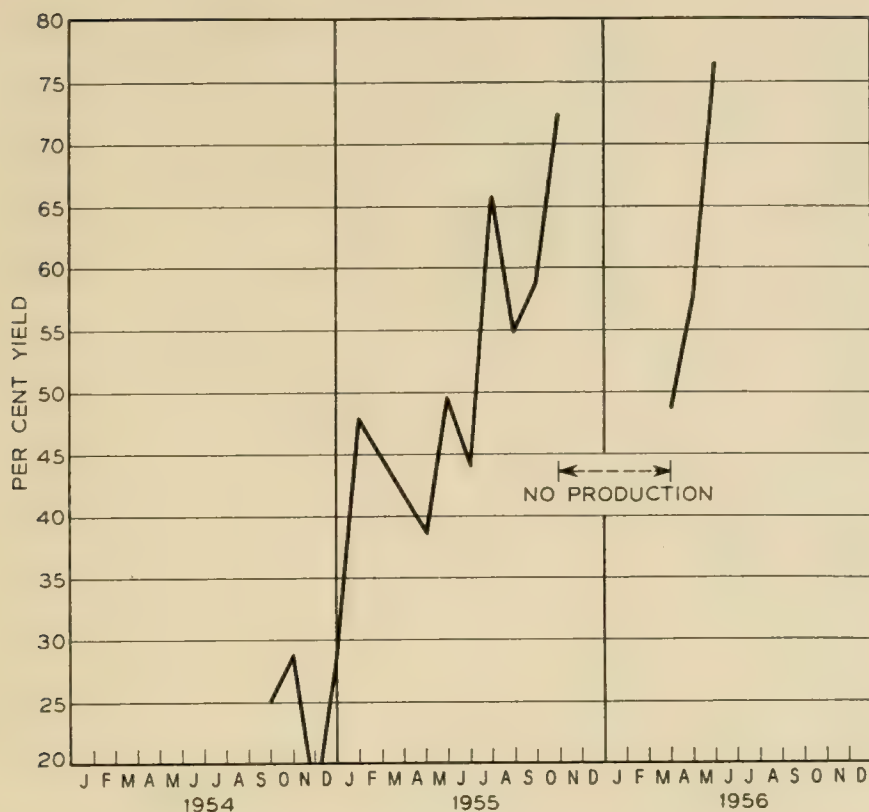


Fig. 6 — Mica capacitor yield.

attachments were necessary to wind these resistors. Other resistors were hand wound on methyl-methacrylate forms, or on the outside of the ceramic containers, for certain types of paper capacitors. Rough adjustments were required of the lengths of resistance wire prior to winding, and close adjustments to resistance values were made after the windings were completed and before leads were attached to resistors. Again it was necessary to provide periodic samples that could be placed on life test by the Laboratories to ascertain that the manufacturing processes were under control. These samples, in all possible cases, were taken from product that would normally be rejected because of some minor defect, but which would not in any way detract from the validity of the life tests. The making of hard solder splices between nichrome resistance wire and gold-plated copper leads, and keeping ceramic parts from coming in contact with metal surfaces and thereby being contaminated because of the ceramic's abrasive characteristics, were two major problems on resistors. Fig. 7 indicates resistor yields.

Inductors comprised 20 different designs, most of which were air core, but there were some for which it was necessary to cement permalloy dust cores into pockets of the methacrylate form, and thereafter using

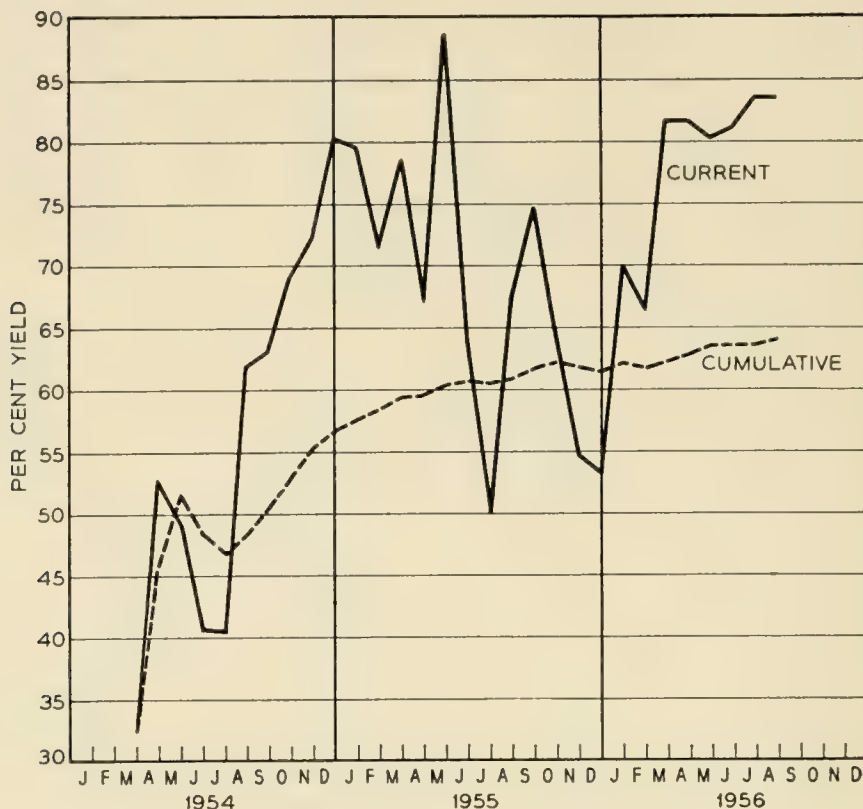


Fig. 7 — Resistor yield.

wire on a shuttle, wind by hand the turns required to produce an inductor. These varied from a very small inductor, smaller in diameter than a pencil, to a fairly large "figure eight" inductor with turns having a major diameter of about $1\frac{1}{4}$ inches. Each layer of a winding was inspected with a microscope to insure that the wire had not been twisted or kinked, or that the insulation was damaged or uneven. Some of the shuttles became fairly long so that they could hold the amount of wire required to make a continuous winding. The operator's handling of this shuttle, as she moved it down around the openings in the methacrylate part, or placed it on a bench to proceed with the interleaving tape, demanded considerable dexterity and concentration to insure that the shuttle was not turned over — which in effect would put a twist in the wire. Although best known means were used to sort cores for their magnetic properties prior to the time a winding was made, the limits on the inductors themselves were so close that subsequently a large number of windings were lost. The best cores that could be selected, plus the best winding practice, could not produce 100 per cent of the inductors within the required limits. Crazing of the insulation on the wire; cementing together of two methacrylate parts or of permalloy cores into pockets of

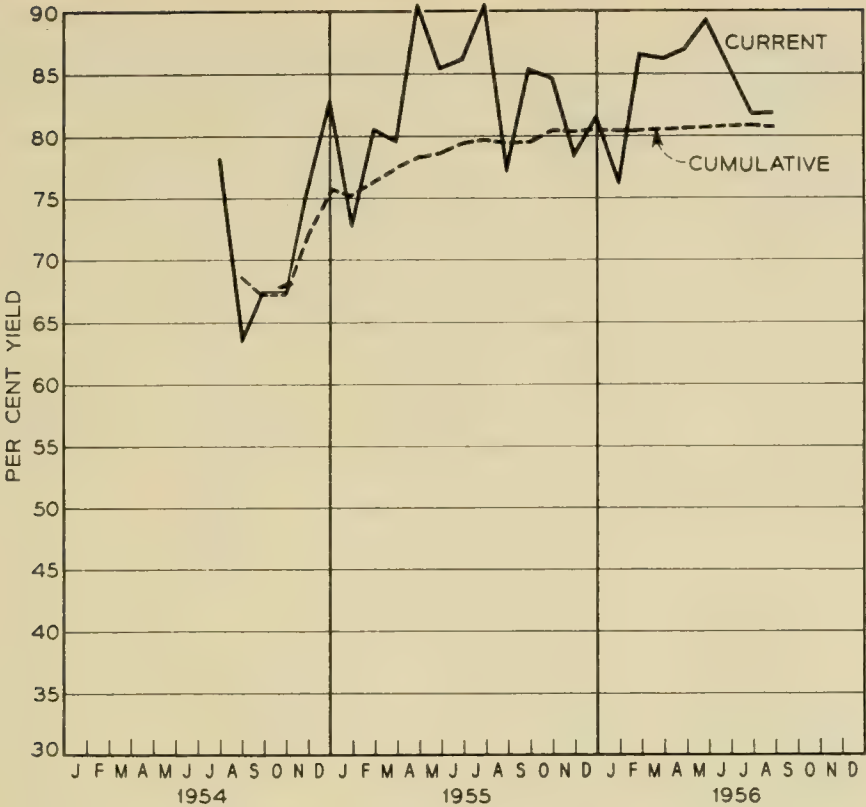


Fig. 8 — Inductor yield.

methacrylate parts, and handling those inductors having long delicate leads, were the most troublesome items on this apparatus. Fig. 8 shows manufacturing yield for inductors.

Networks combined several codes of component apparatus, such as a mica and a paper capacitor, resistor and an inductor. Six networks were used in each repeater unit consisting of two interstage networks, an input, an output, and two beta networks. They demanded a most delicate wiring job in that stranded gold-plated copper wires had to be joined in a small pocket in methyl methacrylate, where a minimum amount of heat can be applied; otherwise the methacrylate is affected. After soldering, a minimum amount of movement of the stranded wire was permitted, inasmuch as the soldered gold-plated copper wire becomes quite brittle.

Repeater Units, are wired assemblies consisting of seventeen sections in which there are six networks, three electron tubes, one gas tube, one crystal, three high voltage capacitors, one dessicator and two terminal sections. The successive build-up of these materials left little chance to make a repair because a splice in a lead was not permissible. It is during this assembly stage that a repeater received its individual identity be-

cause of the frequency of the particular crystal assembled into the unit. A manufacturing yield of 100 per cent was achieved in the assembly and wiring of repeater units.

It was necessary to calibrate the test equipment for this job very closely. Bell Telephone Laboratories and Western Electric worked at length to calibrate the testing details and the test sets for individual networks. Adjustments in components apparatus to bring the network to the fine tolerances required were accomplished by minute scraping of the silvered mica on a mica capacitor or removing turns from wire-wound inductors. The cementing of methacrylate parts, which was a troublesome item on mica capacitors and inductors, also had to be contended with on networks.

PACKING AND SHIPPING COORDINATION

Repeaters were packed in Western Electric specially designed 34-foot long aluminum containers, weighing 1,000 pounds. Forty of these containers were made by an outside firm. Fig. 9 shows two containers tied down in a truck trailer. The repeaters were nested in a pocket of polyethylene bags containing shaped rubberized hair sections in order to cushion the repeaters during their subsequent handling and transportation. The instrumentation required with each case was tested, properly set, and inspected prior to its use on each outgoing case. The instruments were a shock recorder to register shocks in three planes, and a thermometer to register the minimum and maximum temperatures to which the repeater had been exposed. Arrangements were made with a commercial trucking company to provide three specially equipped truck trailers, which could be cooled by dry ice during hot weather and warmed by burning bottled gas during cold weather so as to control temperature within the 20-degree F. to 120-degree F. called for in the repeater specification.

Appointment of a shipping coordinator supervisor added tremendously to the smooth functioning of services and provided the continuing vigilance required to protect repeaters and deliver them to the right place at the right time. His responsibility was to coordinate all the shipping information and arrangements from the time the item was ready for packing at the Hillside plant, through all trucking arrangements to the armoring factory, to the airport, to England, and to follow, with statistical data and reports, each enclosure until we were able to record the date on which the repeater was laid or stored in a depot.

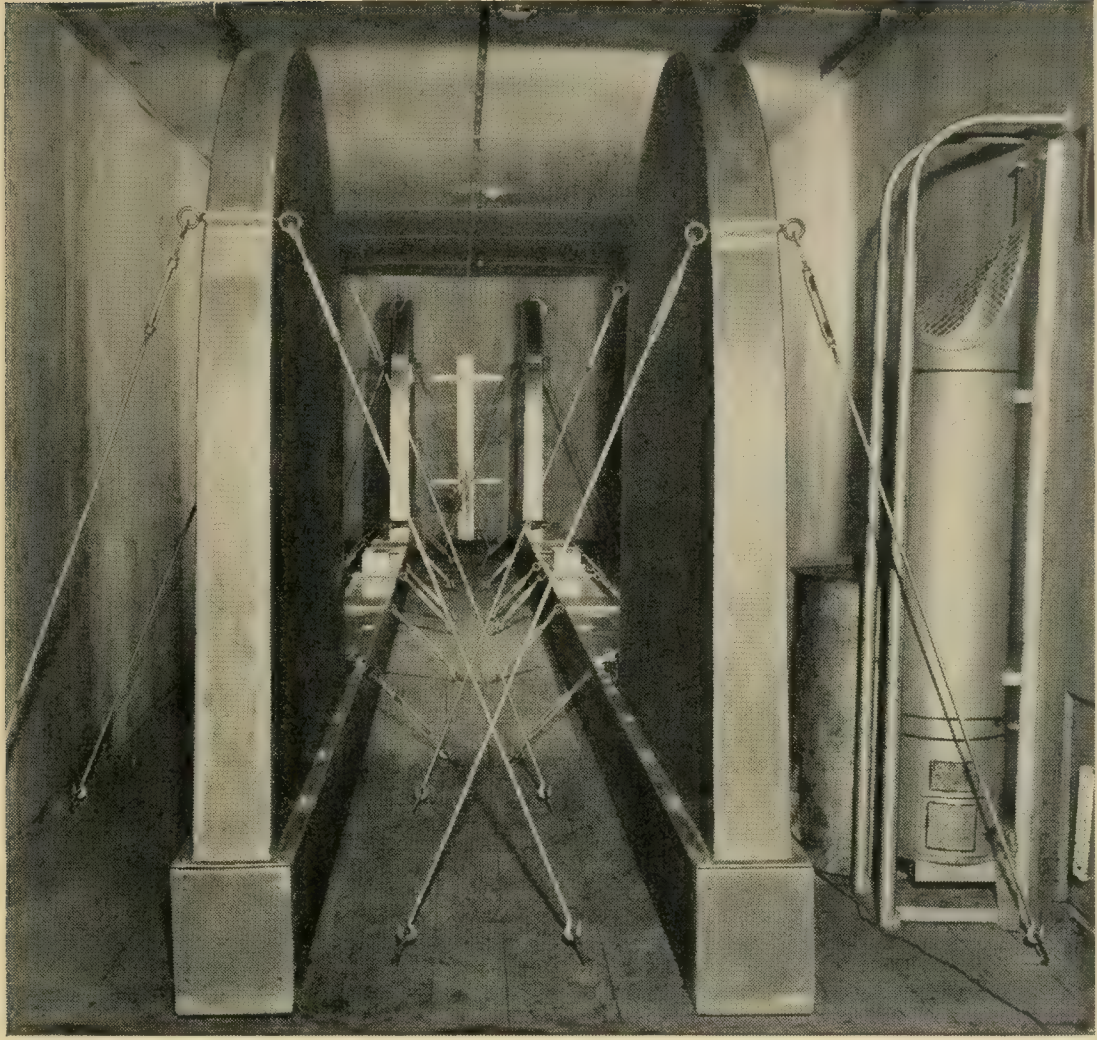


Fig. 9 — Shipping containers.

INSPECTION PLAN AND PROCEDURES

General

It is axiomatic that quality is not obtained by inspection but must be built into the product. However, the Inspection Organization does have the responsibility of certifying that the desired quality exists. Our evaluation indicated that the ordinary inspection "screening" would be inadequate to insure the high degree of integrity demanded and that additional safeguards would have to be provided. These controls were achieved, in a practical way, by:

- (1) Selective placement, intensive training and subsequent qualification testing of all personnel.
- (2) Inspection during manufacturing operations in addition to in-

spection of product after completion, and regulating inspection so that critical characteristics received repetitive examination during the process of manufacture and assembly.

(3) A maintenance program for inspection and testing facilities which provided checks at considerably shorter intervals than is considered normal.

(4) Inspection and operating records and reports that point out areas for corrective measures.

(5) Records of quality accuracy for all inspection personnel as an aid in maintaining the high quality level.

(6) Verification of all data covering process and final inspection as a certification of the accuracy of these data and that the apparatus satisfactorily meets all requirements.

Selection and Training of Inspection Personnel

The quality of a product naturally depends upon the skills, attitude, and integrity of the personnel making and inspecting it. It was realized that in order to develop the high degree of efficiency in the inspection organization necessary to insure the integrity of the product, personnel of very high caliber would be required. These employees would have to be (1) experienced in similar or comparable work, (2) they would have to be precise, accurate and, above all, dependable, (3) in order to reduce the possibility of contamination and damage they would have to be neat and careful, and (4) they would require the ability to work in harmony with other employees, often as a member of a "team," in an environment where their work would be under constant scrutiny.

Most of the inspection employees selected to work at Hillside were transferred from the Kearny Plant and had an average Western Electric service of twelve years. They were hand-picked for the attributes outlined above, and the "screening" was performed by supervision through personal interviews supplemented by occupational tests given by the personnel department. These tests, which are in general use, are designed to evaluate background and physical characteristics, and they were given regardless of whether the employee had or had not previously taken them.

The following group of tests is an example of those given inspectors and testers of apparatus components:

- (1) Electrical — ac-dc theory and application.
- (2) Ortho-Rater — Eye test for phoria, acuity, depth, and color.
- (3) Finger Dexterity — Ability and ease of handling small parts.

(4) Special — Legibility of handwriting, ability to transcribe data and to use algebraic formulae in data computations.

Inspection Plan

The general plan of visual and mechanical inspection consisted of:

(1) Inspection of every operation performed — and in many cases partial operations — during the course of manufacture. This is of particular importance where the quality characteristics are hidden or inaccessible after completion of the operation.

(2) Repeated inspection at subsequent points for omissions, damage and contamination.

(3) Rejection of product at any point where there was failure to obtain inspection or where the results of such inspection had not been recorded.

Most of the visual inspection was performed at the operators' positions to reduce, to a minimum, the amount of handling that could result in damage and contamination.

Visual inspection covered three general categories:

(1) Inspection of work after some or all operations had been completed, such as the machining of parts.

(2) Inspection at those points where successive operations would cover up the work already performed. An example of this is the hand winding of toroidal inductors where each layer of wire was examined under a microscope for such defects as twists, cracks, and crazes in enamel insulation, spacing and overlapping of turns, and contamination before the operator was allowed to proceed with another layer. While being inspected, the work remained in the holding fixture, which was hinged in such a manner as to permit inspection of both top and bottom of the coil. Inductors received an average of 13 and a maximum of 26 visual inspections during winding.

(3) Continuous "over-the-shoulder" inspection, where strict adherence to a process was required or where it was impossible to determine, by subsequent inspection, whether or not specific operations had been performed. In these cases, the inspector checked the setup and facilities, observed to see that the manufacturing layouts were being followed, that the operations were being performed satisfactorily, and that specifications were being met.

ELECTRICAL TESTING

The electrical testing, in itself, was not unusual for carrier apparatus and runs the gamut from dc resistance through capacitance, inductance,

and effective resistance, to transmission characteristics in the frequency band 20–174 kc. What was unusual were the extremely narrow limits imposed and the number and variety of tests involved as compared to those usually specified for commercial counterparts.

The following two examples will serve to illustrate the extreme measures taken to prove the integrity of the product:

(A) One type of Resistor was wound with No. 46 mandrelated nichrome wire to a value of 100,000 ohms plus or minus 0.3 per cent. This resistor received six checks for dc resistance, five for instantaneous stability of resistance and two for distributed capacitance, at various steps in the process which included six days' temperature cycling for mechanical stabilization. This resistor was considered satisfactory, after final analysis of the test results, if: (a) The difference in any two of the six resistance readings did not exceed 0.25 per cent. (b) The change in resistance during cycling was not greater than 0.02 per cent. (c) The "instantaneous stability" (maximum change during 30 seconds) did not vary more than 0.01 per cent. In addition, it was required that the distributed capacitance, minimum 7, maximum 10 mmf, should not differ from any other resistor by more than 2 mmf.

(B) For high voltage paper capacitors, the 0.004-inch thick Kraft paper, which constitutes the dielectric, was selected from the most promising mill lots which the manufacturers had to offer. This selection was based on the results obtained from tests that involve examination for porosity, conducting material and conductivity of water extractions. These tests were followed by the winding and impregnation in Halowax of test capacitors. The test capacitors were then subjected to a direct voltage endurance test at 266 degrees F for 24 hours.

Samples of prospective lots of paper, which have passed the above test, were then used to wind another group of test capacitors that were subsequently impregnated with Aroclor and sealed. 1,500-volt dc was then applied to the capacitors at 203° F for 500 hours. In case of failure, a second sampling was permitted.

After the foregoing tests had been passed, the supplier providing the particular mill lot was authorized to slit the paper. Upon receipt, six special capacitors were wound, using a group of six rolls of the paper being qualified. These capacitors were then impregnated, checked for dielectric strength at 3,000-volt dc, and measured for capacitance and insulation resistance. The capacitors were then given an accelerated life test at 2,000-volt dc, temperature 150° F, for 25 days. Each lot of six satisfactory test capacitors qualified six rolls of paper for use.

Product capacitors were then wound from approved paper, and the dry units checked for dielectric strength at 300-volt dc. Capacitance

was checked and units were then assembled into cans and ceramic covers soldered in place. Assemblies were pressurized with air, through a hole provided for the purpose, while the assembly was immersed in hot water to determine if leaks were present. Capacitors were then baked, vacuum dried, impregnated, pressurized with nitrogen, and sealed off. The completely sealed units were then placed in a vacuum chamber at a temperature of 150° F, 2 mm. mercury, for 3 hours to check for oil leaks. Capacitance was rechecked and insulation resistance measured.

After seven days, capacitors were unsealed to replenish the nitrogen that had been absorbed by the oil, resealed and again vacuum leak tested. An X-ray examination was then made of each individual unit to verify internal mechanical conditions. Capacitors were then placed in a temperature chamber and given the following treatment for one cycle:

16 hours at 150°F; 8 hours at 75°F; 16 hours at 0°F; 8 hours at 75°F.

At the end of ten days, or 5 cycles, the insulation resistance and conductance was measured and a norm established for capacitance.

Capacitors were then recycled for ten days, and, if the capacitance had not changed more than 0.1 per cent, they were satisfactory to place on production life test. If the foregoing conditions had not been met, the capacitors were recycled for periods of ten days until stabilized.

At that time, 10 per cent of the capacitors in every production lot were placed on "Sampling Life Test", which consisted of applying 4,000-volt dc in a temperature of 150°F for 25 days. At the same time, the balance of the capacitors in the lot were placed on production life test at 3,000-volt dc in a temperature of 42°F for 26 weeks. At the end of this time, the insulation resistance was measured and the capacitance checked at 75°F and at 39°F. The difference in capacitance at the two temperatures could not exceed +0.001, -0.005 mf, and the total capacitance could not exceed maximum 0.3726, minimum 0.3674 mf. The capacitance from start to finish of the life test could not have changed more than plus or minus 0.1 per cent.

If all of the preceding requirements had been satisfied, the particular lot of capacitors described was considered satisfactory for use.

The foregoing examples are typical of the procedures evolved for insuring, to the greatest degree possible, the long, trouble-free life of all apparatus used in the repeater.

Radioisotope Test

There were many new and involved tests which were developed and applied to the manufacture of repeaters. One of the most unique is the use of a radioisotope for the detection of leaks under hydraulic pressure.

The initial closure operations consisted of brazing into each end of the repeater housing a Kovar-to-glass seal. These seals are equipped with small diameter nickel tubulations which were used to flush and pressurize the repeaters with nitrogen. After these operations had been performed, one of the tubulations was pinchwelded, overbrazed and coiled down into the seal cavity. The repeater was then placed in a pressure cylinder with the open tubulation extending through and sealed to the test cylinder. A mass spectrometer was then attached to the tubulation and the test cylinder pressurized with helium at 10,000 psi. At the conclusion of this test the repeater was removed from the test cylinder and, after breaking the desiccator diaphragm, the remaining open tubulation was pinchwelded and overbrazed. At this point, it became necessary to determine whether the final pinchweld and overbrazing would leak under pressure.

Since there was no longer any means of access to the inside of the repeater, all testing had to be done from the outside. This was accomplished by filling the glass seal with a solution of radioisotope cesium 134, which was retained by a fixture. The repeater was then placed in a test cylinder and hydraulic pressure applied, which was transmitted to the radioisotope in the fixture. After 60 hours under pressure, the repeater was removed from the cylinder and the seal drained and washed. An examination was then made with a Geiger counter to determine if any of the isotope had entered the final weld.

The washing procedure, after application of the isotope solution, involved some sixty operations with precise timing. In the case of the repeater at the rubber seal stage where both ends were tested, it was desirable that these operations be performed concurrently. This was accomplished by recording the entire process on magnetic tape which, when played back, furnished detailed instructions and exact timing.

RAW MATERIAL INSPECTION

As might be expected, raw materials used in the project were very carefully examined and nothing left to chance. Every individual bar, rod, sheet, tube, bottle or can of materials was given a serial number and a sample taken from each and similarly identified. Each sample was then given a complete chemical and physical analysis before each corresponding piece of material was certified and released for processing. In many cases, the cost of inspection far exceeded the cost of the material. However, the discrepancies revealed and the assurance provided, more than justify the expense.

Detailed records of all raw material inspection were compiled and furnished to the responsible raw material engineer who examined them,

critically, as an additional precaution before the material was released to the shop.

INSPECTION RECORDS

To eliminate, as much as possible, the human element in providing assurance that all prescribed operations had been performed satisfactorily, inspected properly and the results recorded, means were established to compile a complete history of the product concurrent with manufacture. This was accomplished through the provision of permanent data books of semilooseleaf design, which require a special machine for removing or inserting pages.

Each of these books covered a portion of the work involved in producing a piece of apparatus and contained a sequential list of pertinent operations and requirements prescribed in the manufacturing process specifications. Space was provided, adjacent to the recorded information, for both the operator and inspector to affix their initials and the data. A reference page in the front of each book identified the initials with the employees' names. All apparatus was serially numbered and the data were identified accordingly. If a unit was rejected, that serial number was not reused.

These data books, in addition to establishing a complete record of manufacture, provided a definite psychological advantage in that people were naturally more attentive to their work when required to sign for responsibility.

QUALITY ACCURACY

As pointed out previously, every precaution was exercised in selecting and training inspection personnel assigned to the project. However, it was realized at the outset that human beings are not infallible and that insurance, to the greatest degree possible, would have to be provided against the probability of errors in observation and judgment. Quality accuracy evaluation procedures were, therefore, established for determining the accuracy of each inspector's performance.

Quality accuracy checking was performed by a staff of five Inspection Representatives and involved an examination of the work performed by inspectors to determine how accurately it was inspected. Materials which the inspector accepted and those which had been rejected were both examined.

VERIFICATION AND SUMMARY OF DATA

As an added measure of assurance as to the integrity of the product, procedures were established for verifying and summarizing the inspection records for each serially numbered component, up to and including complete repeaters.

Verification involved a complete audit of the inspection records to provide assurance that all process operations were recorded as having been performed satisfactorily, that the prescribed inspections had been made, and that the recorded results indicated that the product met all of the specified requirements. This work was performed by a group of six Inspection Representatives who had considerably experience in all phases of inspection and inspection records.

As the verification of a particular piece of apparatus proceeded, a verification report was prepared which, when completed, contained the most pertinent inspection data, such as:

- (1) Recorded measurements of electrical parameters.
- (2) Values calculated from measurements to determine conformance.
- (3) Confirmation that all process and inspection operations had been verified.
- (4) Identification (code numbers and serial or lot numbers) of materials and components entering into the product at each stage of manufacture.

The verification report usually listed the data for twenty serial numbers of a particular code of apparatus along with the specified requirements. Included, also, was a cross-reference to all the inspection data books involved so that the original data could be located easily. These verification reports were prepared for all apparatus up to and including the finally assembled and tested repeaters.

The following gives an indication of the number of items examined in the verification of one complete repeater:

Items verified in data books	17,593
Items verified on recorder charts	1,142
Calculations verified	1,580
	20,315
Number of entries on verification reports	4,070

Verification reports, in addition to presenting the pertinent recorded data, provided a "field" of twenty sets of measurements from which it was easily possible to spot a questionable variation. For example, it was the adopted practice on this project to examine, critically, any characteristic of a piece of apparatus, in a universe of twenty, which varied considerably from the rest, despite the fact that it was still within limits.

While the number of cases turned up in the verification process which have resulted in rejection of product are relatively few, we believe that the added insurance provided, and the psychological value obtained, considerably outweigh the cost.

Power Feed Equipment for the North Atlantic Link

By G. W. MESZAROS* and H. H. SPENCER*

(Manuscript received September 20, 1956)

Precise regulation of the direct current which provides power for the undersea repeaters in the new transatlantic telephone cable is necessary to maintain proper transmission levels and to assure maximum repeater tube life. The highest possible degree of protection is needed against excessive currents and voltages under a wide variety of possible fault conditions. Furthermore, to minimize the dielectric stresses, a double-ended series-aiding power feed must be used and the balance of these applied voltages must be maintained in spite of substantial earth potentials. This paper describes the design features which were employed to attain these objectives simultaneously, while eliminating, for all practical purposes, any possibility of even a brief system outage due to power failure.

INTRODUCTION

The principal objectives in the power plant design for the Transatlantic cable system were as follows:

1. To stress reliability in order to guarantee continuous dc power to the electron tubes that form an integral part of the submerged repeaters. This is essential, not only to be able to maintain continuous service, but to prevent cooling and contraction of the repeater components, especially the tubes.

2. To provide close dc cable current control to ensure constant cathode temperature and regulated plate and screen potentials for the repeater tubes. These operating conditions are essential both for obtaining maximum life from these tubes and for maintaining constant transmission level.

3. To control and limit the applied dc cable potentials in order to minimize the dielectric stresses. The life of certain capacitors in the repeaters is critically dependent upon these stresses. Moreover, momen-

* Bell Telephone Laboratories.

tary high potentials increase the chances of corona formation and insulation breakdown.

4. To protect the cable repeaters from the excessive potentials or currents to which they might be subjected after an accidental open or short circuit in the cable.

5. To compensate for earth potentials up to 1,000 volts, of either polarity, that may develop between the grounds at Oban and Clarendville during the magnetic storms accompanying the appearance of sun spots and the aurora borealis.

6. To provide adequate alarms and automatic safety features to ensure safe current and voltage conditions to both the cable and the operating personnel.

DESIGN REQUIREMENTS

Reliable Cable Power

The first basic problem of design was to select a reliable source of dc power for energizing the cable repeaters. Although a string of batteries, on continuous charge, is perhaps the most dependable source of direct current, such an arrangement is not attractive here. A complex set of high-potential switches would be required for removing sections of batteries for maintenance and replacement purposes. Protection of the repeater tubes from damage during a cable short circuit would be difficult. Facilities to accommodate changing earth potentials would be cumbersome. Furthermore, the problem of hazards to personnel would be serious.

The use of commercial ac power with transformers and rectifiers to convert to high potential dc would expose the cable to power interruptions even with a standby diesel-driven alternator, because of the time required to get the engine started. A diesel plant could be operated on a continuous basis, but this prime power source would also present a considerable failure hazard even with the best of maintenance care. The two-motor alternator set, used so successfully in the Bell System's type "L" carrier telephone system, was adopted as representing the most reliable continuous power source available. This set normally operates on commercial ac power, but when this fails, the directly-coupled battery-operated dc motor quickly and automatically takes over the drive from the induction motor, to prevent interruption of the alternator output. Here the storage battery is still the foundation for continuity, but at a more reasonable voltage.

As described later, the possibility of a system outage resulting from fail-

ure of this two-motor alternator set has been essentially eliminated by using two such sets, cross-connected to the rectifiers supplying power to the two cables, with a continuously operating spare for each set, automatically switched in upon failure of the regular set.

The regulating features of the rectifiers will be described in a later section. In the present discussion of reliability it is sufficient to note that series regulating tubes are used, which are capable of acting as high-speed switches, through which two rectifiers can be paralleled. Thus either rectifier can accept instantaneously the entire load presented by the cable. In each regulator the series tubes carrying the cable current are furnished in duplicate and connected in parallel to share the cable load, a single tube being capable of carrying the entire load. These current regulators are operated from separate ac sources to protect against loss of cable power because of failure of one of the sources of ac power.

Cable Potentials

To minimize the cable potentials, half of the dc power is supplied at each end of each cable, the supplies being connected in series aiding. With this arrangement, as shown in Fig. 1, the dc cable potential at one end of each cable is positive with respect to ground while at the other end the potential is negative. This places the maximum potential and risk on the repeaters near the shore ends, which are more readily retrieved, while the repeaters in the middle of the cable, in deeper water, have potentials very near to ground. The power equipment would be simpler with a single-ended arrangement, but at the penalty of doubling the dielectric stresses in the entire system, which would be prohibitive. A balanced power feed could have been attained at the expense of power separation filters in the middle of the cable or a shunt impedance of appropriate size at the midpoint. The resulting complications, in-

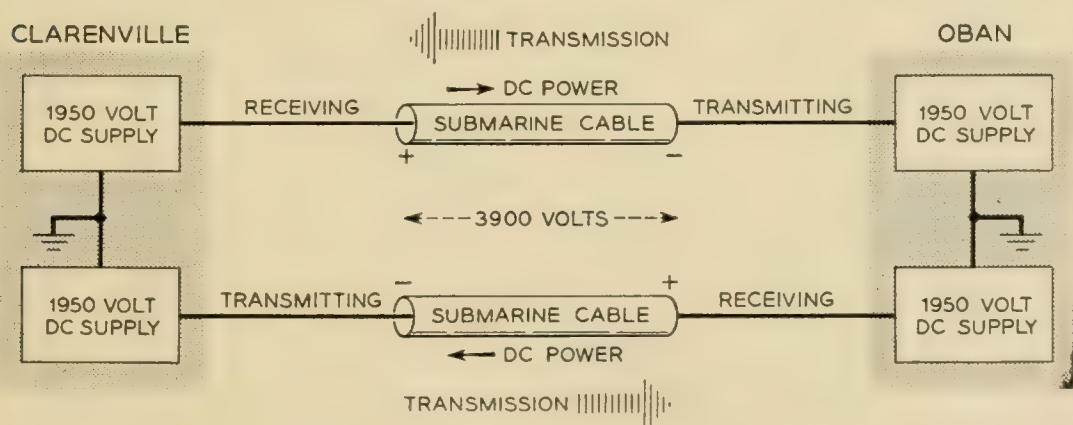


Fig. 1 — Cable voltage supply.

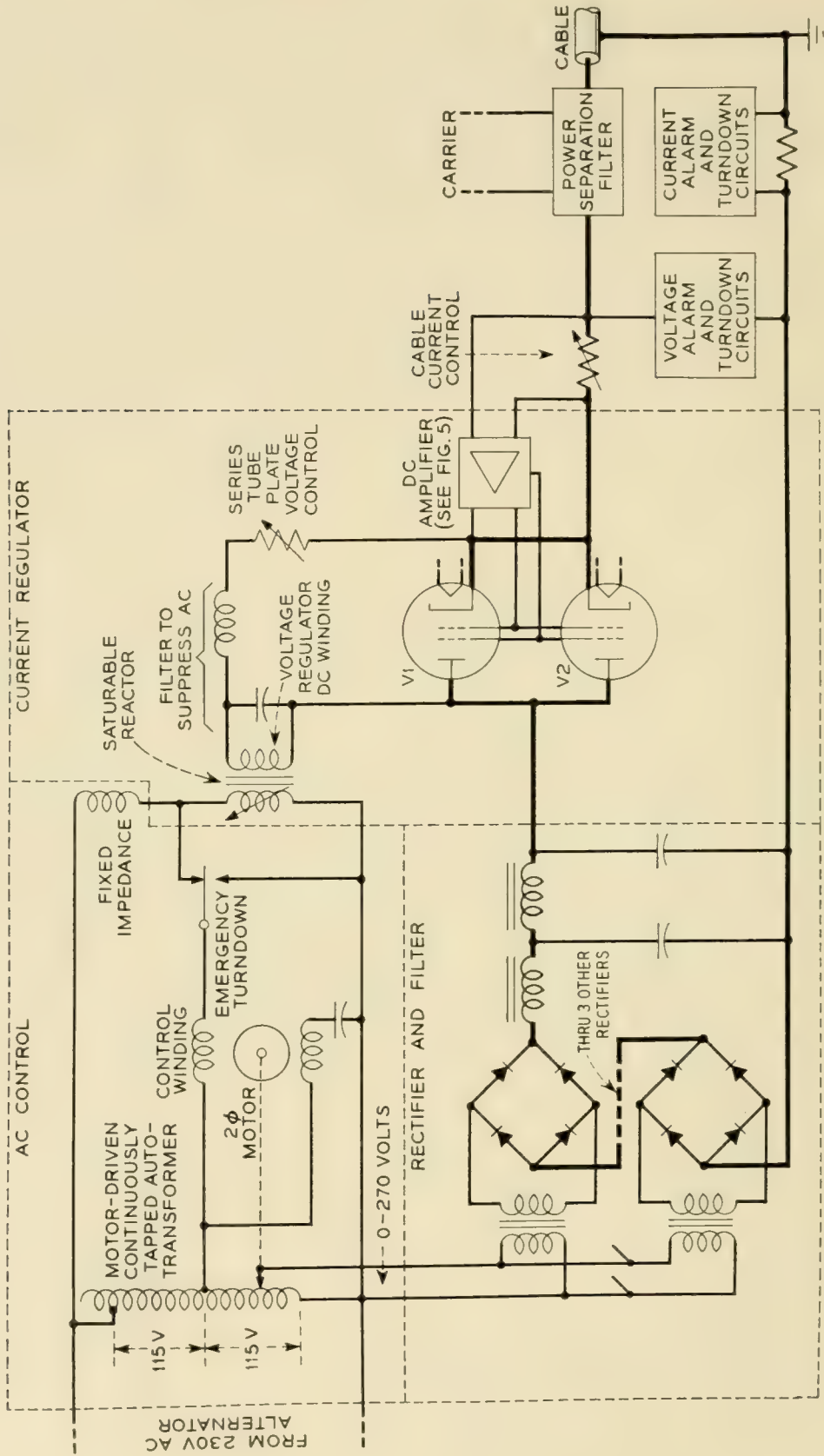


Fig. 2 — Simplified circuit of the regulating system.

cluding difficulty in the location of a cable fault, could not be justified for the sake of simplification of power-plant design and operation.

• The requirement that minimum cable potentials be maintained during and after severe earth potential disturbances necessitates variable output voltages from the supplies at both ends, and this introduces problems in continuous voltage balance and regulation stability. The design features which yield the required performance are described in a later section.

DC Cable Current Regulation

The salient requirements in performance of the constant current regulator are listed below:

a. The regulator must have extremely fast response to hold the cable current within a few milliamperes of its nominal value should a short circuit develop in the cable. Thus damage to the heaters of the repeater tubes, as well as excessive induced transient voltages in the repeater transformers is avoided. The probability of a short circuit is higher near the shore ends where the water is shallow and sea traffic a factor. The regulator must be capable of absorbing the reduction in power to the cable, while maintaining current control under normal conditions. This sudden exchange in power from cable to regulator may be as much as 2,000 volts at 0.25 ampere.

b. The cable current should be maintained constant within 0.2 per cent of its nominal value for normal variations in ac supply, gradual earth potential changes, and ambient temperature changes. This degree of regulation allows an adequate safety factor in maintaining a constant transmission level.¹

c. The regulators, in conjunction with the power separation filters and the rectifier filters, must limit the power supply noise at the cable terminals to a peak-to-peak value less than 0.02 per cent of the dc supply potential.

d. The cable current must be adjustable over a range of 225 to 245 milliamperes to compensate for repeater tube aging.²

e. The regulators must operate in parallel in such a way as to ensure continuity of power should one fail or be removed from service for maintenance. This of course implies that regulators can be switched in and out of service without causing surges in the cable current or voltage.

f. The series-aiding arrangement, with rectifiers at each end of the same cable, must be stable.

g. The regulators should be capable of being serviced at low poten-

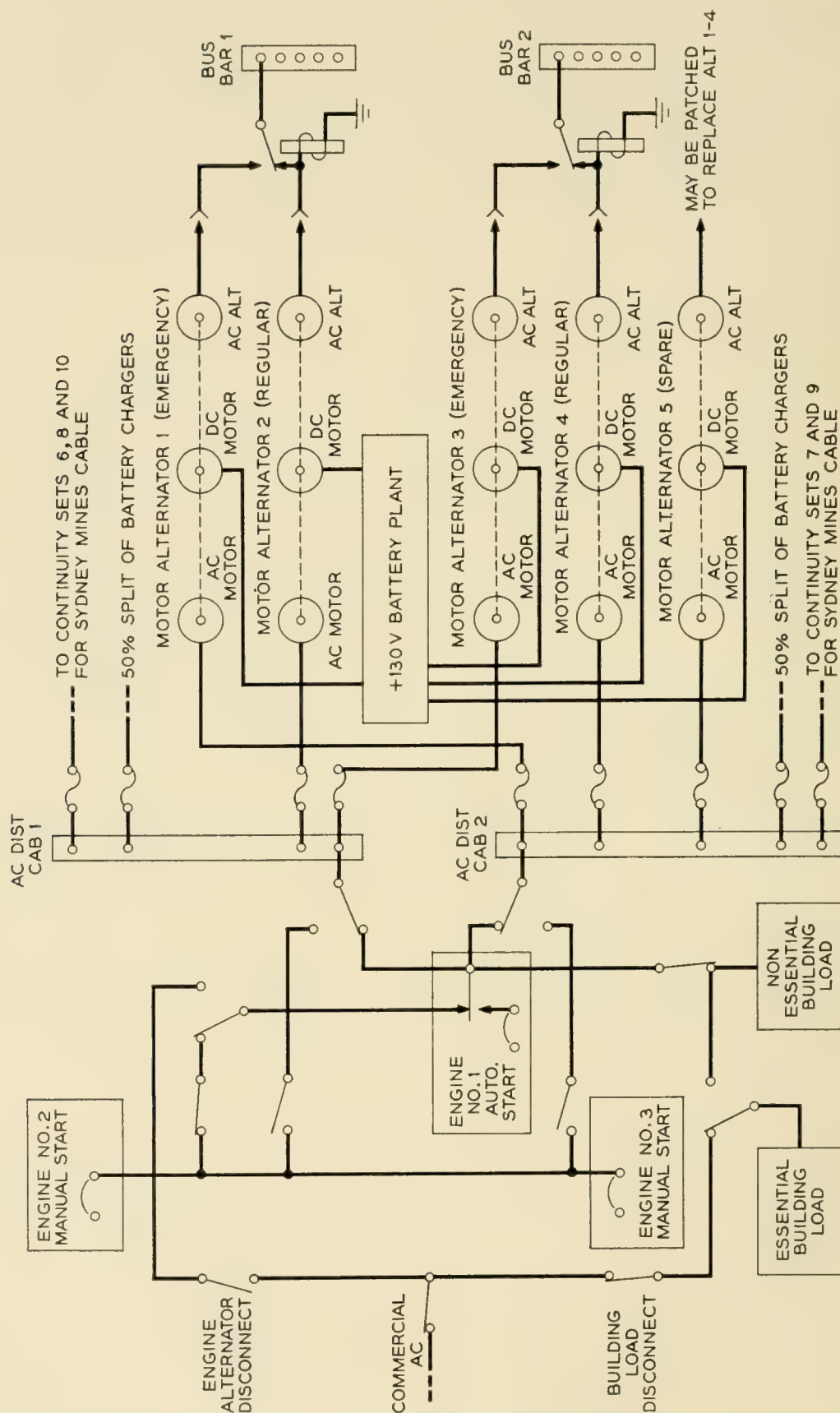


Fig. 3 — Continuous ac power.

tials, when in the test position, in order to protect maintenance personnel.

h. The current regulator should be of the fail-safe type so that impairment of any of the regulator components will not permit excessive rise in cable current. In the event of component trouble, an aural or visual alarm should occur.

It was decided that a high speed electronic constant-current regulator backed up by a slower speed servo system, as shown in Fig. 2 and discussed in detail later, would best meet the above requirements. In this way, fast response with high gain is combined with wide regulating range, yet the efficiency is high and the load-handling capacities of the various components are held to a minimum. With regard to simpler alternatives, the electromechanical type of current regulator, using relays and a motor-driven rheostat, is too slow to protect the repeater tubes from a cable short circuit and its accuracy is insufficient to meet the regulation requirements. The all-magnetic type of regulator is possibly most dependable but it does not readily provide either the speed of response or the wide regulating range needed.

GENERAL DESCRIPTION

Prime and Standby Power Source

Commercial service is considered the normal prime source of power for the cable, although at the Clarendville terminal commercial power was not available at the time of installation. Anticipating this condition, a reserve plant consisting of three 60-kw diesel alternators was installed and the distribution circuits were arranged, as shown in Fig. 3, to provide partial or total use of the commercial service. Initially all cable power was supplied by diesel operation, alternating the prime movers on a weekly basis. These sets are paralleled manually when they are interchanged, to prevent an interruption in the 60-cycle supply. It may be noted that Engine No. 1 is arranged as an automatic standby whether prime power is provided by diesels or by commercial service.

The switching and distribution arrangements are designed to be essentially failure-proof. At Clarendville, for example, two ac distribution cabinets, each capable of being fed from two sources, were provided in separate locations. The normal source through Engine No. 1 control bay can be readily by-passed directly to the manual diesels, should Engine No. 1 control bay be disabled. Furthermore, allocation of charging rectifiers, control circuits, ac motors for continuity sets, etc., has been

made in such a manner that loss of one cabinet alone will have minimum effect on the cable power supplies or office loads. At Oban, where 50-cycle commercial service is normally used, special distribution arrangements have been provided to give maximum power supply reliability, with three manually operated 50-cycle, 90-kw, diesel-alternator sets arranged for standby service. The diesels at this terminal are larger to care for greater local power loads.

Continuous AC Power From Two-Motor Alternators

At both cable terminals, two reliable ac buses supply power to the dc cable regulating bays. Each of these buses is fed from a continuously operated, self-excited, single-phase, 230-volt alternator normally driven by a 3-phase induction motor on the same shaft with a 130-volt dc motor. Each regular alternator is backed-up by a similar emergency alternator running at no load. A fifth motor-alternator is provided which can be used whenever any other set is out of service for routine maintenance or repair.

As alternator loads are essentially constant, and since induction motor speeds are fairly insensitive to power supply voltage variations, alternator outputs are set by fixed adjustments of their field rheostats. Supply voltages are monitored to control automatic transfer to dc motor drive whenever the supply voltage drops below 80 per cent of the normal value. Fig. 4 shows the normal running circuit for an alternator set, with the dc motor connected to the battery through a resistance of 75 ohms inserted in the armature circuit. The field resistance FR is preset so that when the battery is driving the set, the speed matches that of

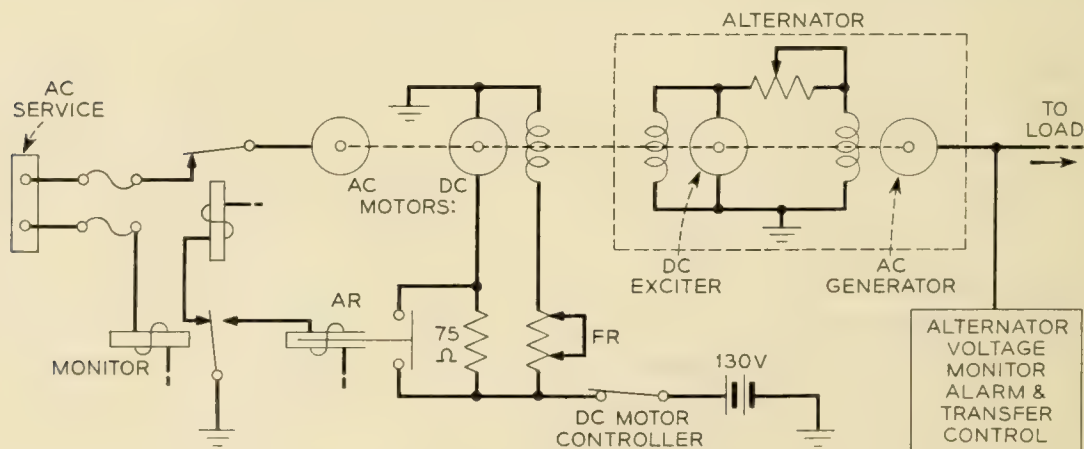


Fig. 4 — Two-motor alternator set.

the ac drive when the battery voltage is at the mean discharge value. During ac drive, the EMF generated in the dc motor armature is a few volts below that of the battery. Accordingly, when the fast acting contactor AR shorts the 75-ohm resistance in the armature circuit, the motor finds itself essentially at the desired operating flux condition and a smooth pickup of drive occurs. Oscillograms indicate that the interval between failure of ac power and operation of contactor AR is less than 0.1 second. During the transfer from ac to dc drive, the change in the nominally 230-volt output is less than 5 volts.

Return to ac drive is delayed approximately 20 seconds after the ac supply voltage has returned to normal to allow time for the ac to stabilize. Fixed field settings for both dc motor and alternator fields provide simple control arrangements without the overspeed or overvoltage hazards which automatic regulators might add. Alternator output is monitored, however, to give alarms for voltage changes exceeding ± 5 per cent and to control transfer to dc drive if the output should drop more than 10 per cent for any reason. When the latter occurs, the machine locks on dc drive. This feature guards against ac motor failure or low ac drive speed because of low supply frequency without low supply voltage.

Failure of the alternator output after transfer to dc drive causes the set to stop and automatically transfer the load to its emergency alternator. This transfer causes a break in the alternator supply to its bus, but cable power is maintained constant by the parallel dc regulating bay fed from the other alternator bus.

Battery Plants and Distribution

Battery power for dc motor drive is supplied from a 66-cell, 1,680-ampere-hour battery at Clarendville and from two 68-cell, 1,680-ampere-hour batteries at Oban. The latter station has double capacity to provide stand-by power for the ac supplies to the inland transmission equipment.

At both cable terminals, control battery for the small alternator plants and the cable dc regulating equipment provides 24 volts and is split so that a fuse or a battery failure on either supply will not interrupt cable power. To guard against so remote a hazard as loss of a common battery for this vital control, two separate 24-volt power plants have been provided with one half of the critical control circuits furnished from each plant.

In addition to supplying dc motor power, the 130-volt battery at Clarendville supplies current to the carrier terminal and test equipment

through voltage dropping resistors which are normally in the circuit to hold the load voltages below a maximum of 135 volts. These resistors are automatically shorted to maintain a minimum load voltage of 125 volts when the battery voltage drops. A voltage detecting relay also steps the fixed field adjustment of each dc motor when the battery nears its final discharge voltage so that dc motor speed is kept within about ± 3 per cent of normal ac drive speed during battery operation. The 66-cell battery is floated at 143 volts by means of voltage regulated rectifiers and, after a discharge, is recharged by automatic operation of a regulated 100-ampere motor-generator set. As indicated in Fig. 3, the rectifiers for this plant are connected so that loss of one service cabinet will still leave sufficient charging capacity to float the load from the other cabinet.

Rectifiers and Associated Controls

As mentioned earlier, each cable is supplied at all times by two regulating bays in parallel, each operating from a separate ac source and each capable of taking over the cable load should the other fail. Thus failure of an alternator supply or of a regulating bay itself will not interrupt the cable power. A spare regulating bay for each cable is arranged for replacing either of the two regular bays and may be connected in parallel with the other two without overloading the associated 2.5-kva alternator.

As shown schematically in Fig. 2, the ac supply to the rectifiers is controlled by a continuously-tapped variable autotransformer, operated normally by a two-phase low-inertia reversible motor, with provision for manual adjustments also. The autotransformer output is stepped up and rectified by a series arrangement of five rectifiers, each to give a maximum of 550 volts for a total of 2,750 volts when the autotransformer is at the upper limit. The high voltage rectifier output is filtered and supplied to the cable through the current regulating unit, cable connecting switch, and common cable control circuit. The cable current is regulated by controlling the plate to cathode drop across the electron tubes through which this current flows. A dc amplifier, which derives its signal from a resistance in series with the total cable current, varies the grid bias of the series regulating tubes. The series tubes have a control-drop range from about 150 volts to the full supply voltage, the dc amplifier being capable of driving the regulating tubes to cut-off. However, to protect the series tubes and to keep them operating at practically a constant plate voltage of 300 volts, a servo system automatically raises and lowers the output from the autotransformer by means of the motor-

driven ac control unit whenever the series tube plate voltage varies more than 25 volts from the normal value. For example, a 2 per cent change in input ac would change the rectifier output of 2,300 volts (corresponding to a cable supply voltage of 2,000 volts) by 46 volts, which would increase the series tube drop to 346 volts. This increase in voltage would raise the signal current through the control winding of a saturable reactor which forms one leg of a balanced bridge in the ac motor control circuit, and thus cause the autotransformer to be driven down, lowering the rectifier output until the series tube drop is restored again to approximately 300 volts.

Overvoltage and Overcurrent Protection and Alarms

While the power for the cable is electronically regulated, protective features are provided to guard against abnormal cable current or voltage. The first order of protection is a ± 2 per cent cable current alarm given by a voltage relay which operates from the drop across a resistor in the ground return circuit to the dc bays. A second voltage relay, set for 5 per cent high cable current, also monitors the current in the ground return side and in conjunction with the current-monitors in the high voltage side, limits the cable current by operating the motor-driven autotransformers until the current is within 5 per cent of normal. The voltage of the ungrounded side, however, is much too high for direct connected voltage or current relays. Therefore, magnetic amplifiers have been used to obtain isolated metering of the cable current. Two of these devices measure the current in the ungrounded side of the common power supply lead to the cables. Of the three current-monitors available, two must operate before turndown functions, to prevent false turndown because of faulty metering.

Voltage protection is provided by means of magnetic amplifiers in shunt across the common power supply to the cable. Here, three monitors are arranged so that any two can reduce power, by means of the turndown control, if the voltage rises to the maximum allowable value for which the voltage relays are adjusted. These monitors draw about 1.5 milliamperes each, but are connected on the supply side of the cable-regulating resistance so as not to affect cable regulation. They guard against excessive voltage resulting from an open circuit where voltages around 4,000 volts could otherwise occur. They also guard against high voltage caused by earth potentials or unbalance between voltages at opposite ends of the cable. When set for a ceiling voltage of 2,600 volts, a maximum of 3,000 volts occurs on open circuit on the first rise, after

which the voltage holds within $2,600 \pm 200$ volts as the turndown relays operate and release to maintain the ceiling voltage. A fourth magnetic-amplifier voltage detecting relay provides an alarm for ± 5 per cent excursions in cable voltage from the normal value. Other alarms are provided to indicate low output in either of the two parallel regulating bays, relay troubles, loss of magnetic amplifier ac control voltage, and fuse failures.

To limit the rate of change in the cable current under short circuit conditions and to reduce the rise in voltage at the repeaters on open circuit failure, an inductance of about 36 henries is connected in series with the cable circuit, and physically close to the cable termination, so that any failure in the power supply would have the advantage of this surge-limiting element.

Metering

Metering of the cable current is a very important part of the power plant design. Not only are the cable current ammeters needed to set the value of current desired, but their ability to indicate absolute current values assists in obtaining stable regulation between the two ends of the cable. The meters provided for this purpose are suppressed-zero, magnetically and statically shielded 150–300 milliamperes, large scale ammeters, with 0.5 per cent accuracy. One of these meters is connected in the ungrounded side and another in the grounded side of the cable supply circuit to provide an accuracy check and to indicate any ground leakage current in the supply circuit. They are connected to highly accurate 1-ohm four-terminal shielded resistors acting as shunts in the cable current circuits with their shunt leads arranged for switching to a calibration box for checking accuracy and for adjustment. This box employs a Weston laboratory standard cell, essentially a single-point potentiometer with the usual galvanometer, acting as a calibration standard at the 225-milliamperes point. Meters calibrated at Oban for 225 milliamperes were expected to be within 0.2 per cent or 0.5 milliamperes of those similarly calibrated at Clarendville, and at present are within about 0.2 milliamperes.

The cable current is also indicated by a recording ammeter. Meters in each regulating bay indicate the division of current between paralleled regulators. These meters have only 1 per cent accuracy but are satisfactory for adjusting load balance between parallel bays and also are used in turnup of power on a particular bay.

Cable voltage is read on a large scale voltmeter reading 0–3,000 volts and having ± 0.25 per cent accuracy. Since the accuracy is not critical,

this meter is not arranged for calibration. However, the series resistors incorporate 6,000-volt components for an extra degree of reliability. The voltmeter with its series resistor is normally connected across the common cable power supply ahead of the cable current regulating point. It can, however, be switched to read the cable voltage nearer the cable termination. In the latter position, it reduces the cable current by about 1 milliamperes, causing an unbalance in cable regulation, and therefore is not normally left in this position. The cable supply voltage is also indicated by a recording voltmeter.

Other meters are provided to indicate series tube plate voltages, dc rectifier voltages, ac input voltages, series tube currents, test currents for adjusting mag-amp operating limits, and the difference in current between the positive and negative power supplies to ground. Since many of these instruments operate at high potential, a special design was used with the operating mechanism and scale depressed in the instrument case approximately 1 inch. This both eliminated possible electrostatic effects on the instrument pointers and served as a safety measure.

DC REGULATION

DC Amplifier

The direct-coupled two-stage amplifier shown in Fig. 5 is characterized by potentiometer coupling and cold-cathode gas-tube voltage stabilizers. The biases are selected high enough so that linear operation is assured even for a short circuit at the cable terminal. As is apt to be the case in a direct-coupled amplifier, cathode temperature in the low-level stage is critical. In this amplifier, a 5 per cent change in heater voltage results in a 0.5 milliamperes change in the cable current.

The required precision of current regulation in these power supplies can be expressed as representing a source impedance of not less than 100,000 ohms. To meet this requirement, the gain of the dc amplifier was made as large as practicable with the plate and screen potentials available from gas-tube regulators. Interstage network impedances are high to reduce the shunt losses, and are proportioned to provide the large biases mentioned above. Gain adjustment is provided by a variable resistance in series with the cathode of the first stage.

Shaping of the loop gain and phase characteristics to obtain margins for stable operation is accomplished by means of the RC shunt (R_9C_1) across the plate resistor for the first stage. The amplifier gain and phase characteristics without this compensation are shown in Fig. 6. These

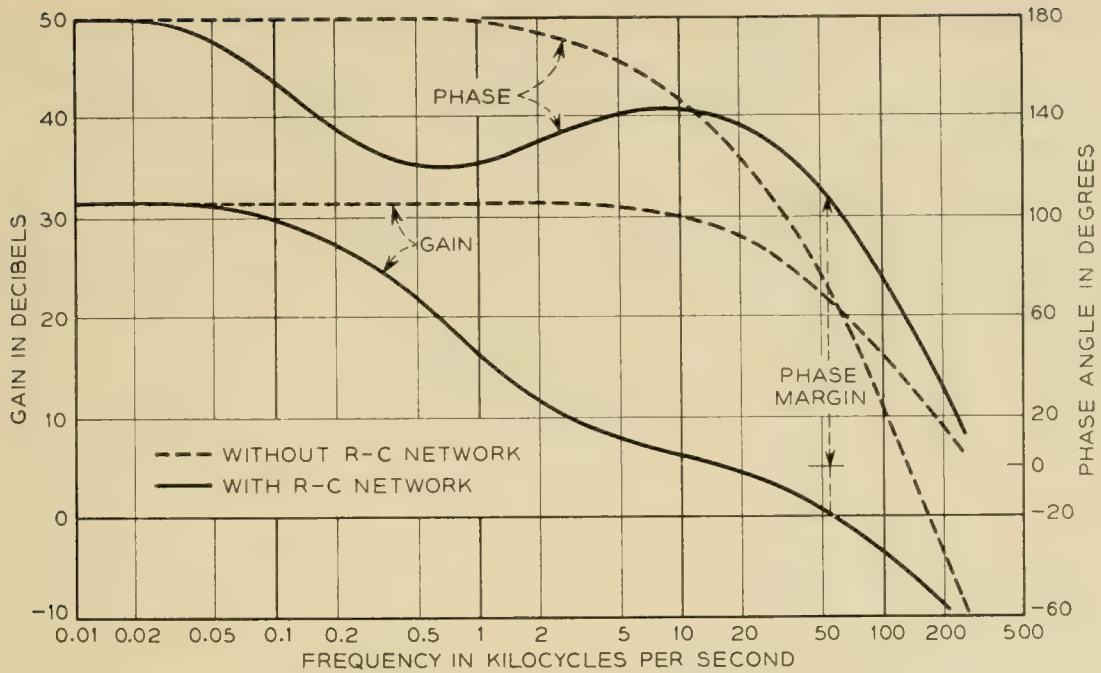


Fig. 6 — DC amplifier gain and phase, experimental model.

data were obtained by opening the feedback loop at the control grid of the first stage, applying normal dc bias plus a variable frequency ac signal to the grid of the tube, and measuring the magnitude and relative phase of the return signal.

The corresponding characteristics with the compensating network in place are also shown in Fig. 6. The compensating network effectively puts a relatively low-impedance shunt across the interstage network at the higher frequencies, resulting in a "step" in the gain characteristic. A secondary effect is the phase shift in the transition region. The calculated "corner frequencies" are 2,800 and 195 cps, respectively, chosen on the basis of the criteria (1) little effect on regulator gain at 100 or 120 cps, the most prominent rectifier ripple frequency, and (2) a gain step of something above 20 db with no appreciable contribution to the phase shift at frequencies above 30 kc. The calculated loss at 120 cps is 1.2 db with a maximum phase shift of about 60 degrees at the median frequency. These results agree quite well with the measured data plotted in Fig. 6.

As indicated in Fig. 6, the phase margin at the gain crossover frequency of 55 kc was somewhat over 100 degrees for the experimental model on which these measurements were made. The gain margin could not be measured readily but is clearly substantial. On production units, larger wire sizes and longer lead lengths resulted in lesser, but still satisfactory stability margins, as shown in Fig. 7, the phase margin being somewhat

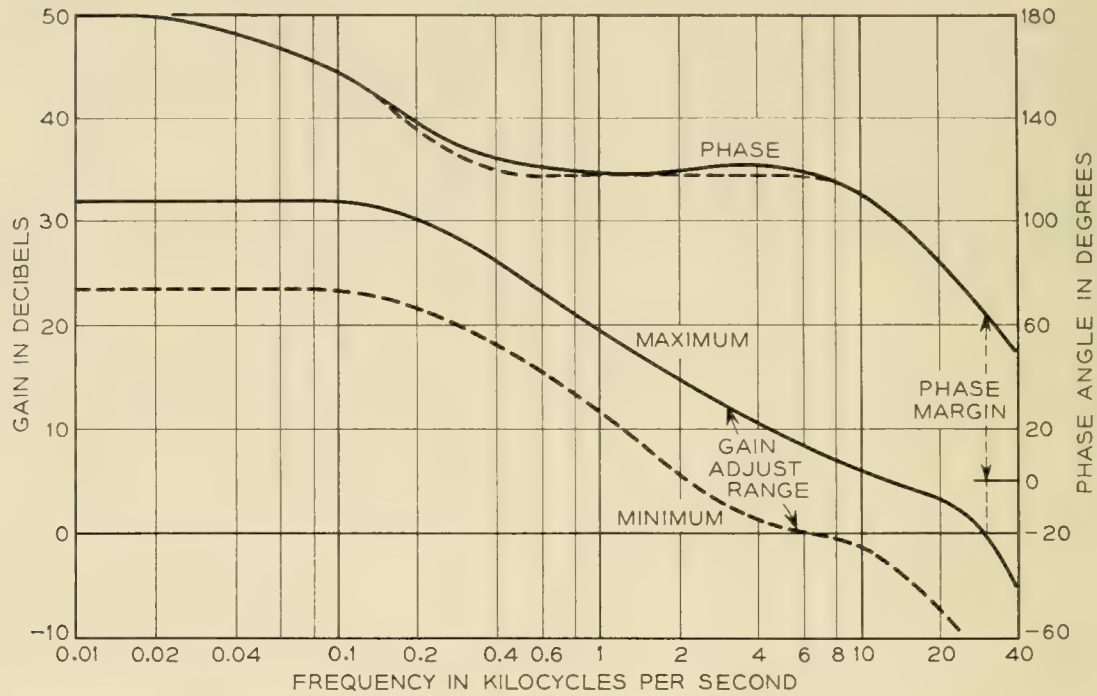


Fig. 7 — DC amplifier gain and phase, production model.

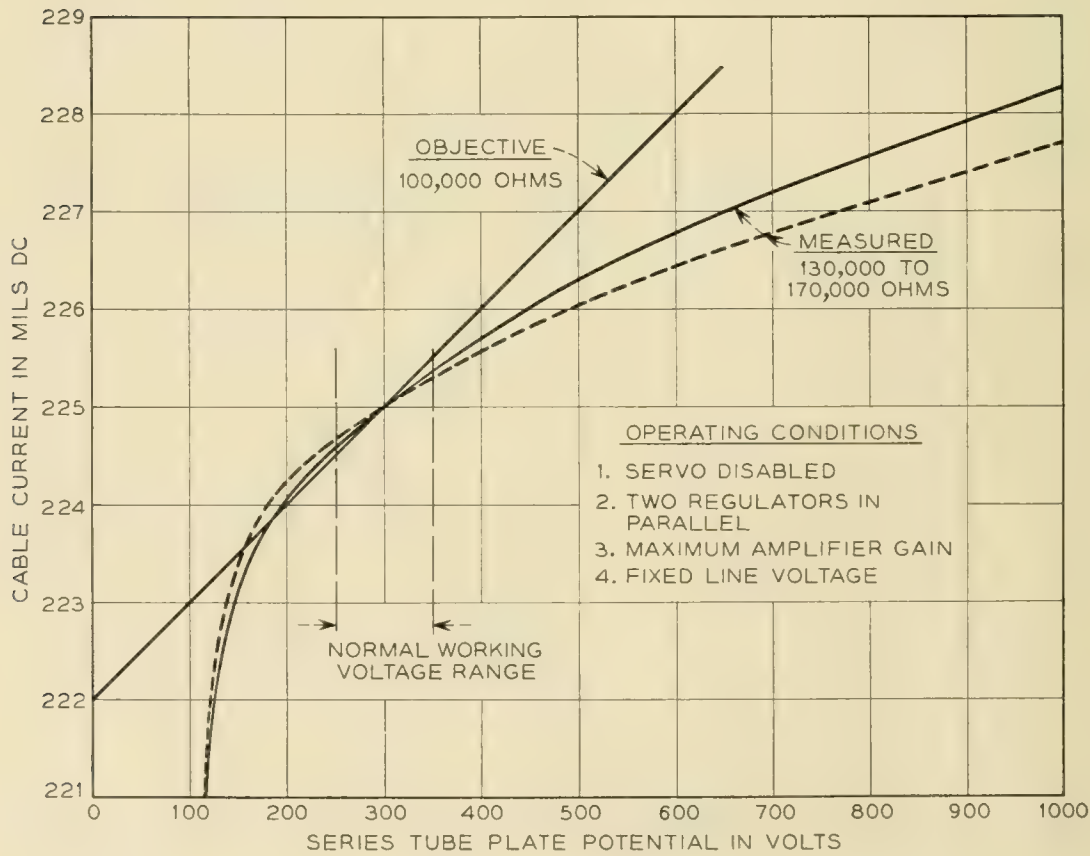


Fig. 8 — Load regulation.

over 60 degrees. Fig. 7 also shows the characteristics at the extremes of gain control, the range of control being about 8 db.

Fig. 8 shows the measured performance of the dc regulators, the servo system being disabled in order to obtain a plot of the performance of the dc amplifier and associated circuits. Twenty-two regulator units were manufactured and measured and the curves of Fig. 8 show the extreme limits observed, the differences between individual regulators being due primarily to differences between electron tubes. The measured range of source impedance, 130,000 to 170,000 ohms, allows margin for regulator tube aging above the 100,000-ohm objective.

AC Servomechanism

As noted earlier, the servo system shown in Fig. 2 is part of the current regulating scheme and holds the series tube plate potential within reasonable limits by adjusting the rectifier input voltage. In an emergency, a "turndown" feature, operated from several remote points, either manually or automatically, will reduce the autotransformer output to zero in less than two seconds. For simplicity, only the manual turndown feature is shown in Fig. 2. It operates simply by switching one end of the motor control winding from one corner of the bridge to the other, thus applying half of the input voltage to the control winding.

Manual operation of the autotransformer tap is provided to raise the cable current slowly, either initially or after a turndown. In manual operation a dynamic brake, consisting of a short circuit on the motor control winding, prevents the motor from creeping or coasting when the operator releases the handwheel, as it otherwise would since the fixed phase of the two-phase motor is always energized. The turndown feature takes precedence over the short circuit of the motor control winding, automatically, to energize the motor should the operator inadvertently cause abnormally high cable voltage or current.

One essential feature of the servo design is the dead band of the series tube plate voltage in which the servo remains stationary, even though there are small changes in the incoming signal. This band can be varied from 10 to 100 volts under control of a gain-adjust potentiometer across the control winding of the two-phase motor. Without this dead band, the servo would be constantly in operation correcting for small random variations in line voltage or earth potentials. Furthermore, since it is extremely difficult to set the current regulators at the two ends of a cable to exactly the same current, the servo dead band permits some margin of error. Otherwise the servo associated with the current regu-

lator trying to regulate for a slightly higher cable current, would drive its rectifier voltage to its stop or maximum output, unbalancing the cable voltages.

*System Stability**

A complete analysis of system stability represents an exceedingly formidable, if not impossible, task. It has been established analytically that for a linear network the two dc regulators in parallel and the system as a whole are unconditionally stable. The details of this proof are too long to be presented here but the line of reasoning with respect to the overall system is as follows. The system of Fig. 1 is symmetrical about a vertical plane through the middle of the figure. Under these conditions, the system will be stable if, and only if, the following three simpler systems† are stable:

- (a) A power supply short-circuited;
- (b) A power supply feeding an impedance equal to twice that of the half cable short-circuited; and
- (c) A power supply feeding an impedance equal to twice that of the half cable open-circuited.

The transfer function of the servomechanism was measured over the frequency range of principal interest, 0 to 1 cps, the behavior near zero frequency being determined from the asymptotic slope of the unit step response.‡ In this frequency range the dc amplifier gain is a real constant, flat gain and negligible delay as previously shown, therefore only the ac servo feedback loop characteristic has to be known to predict the stability of condition (a). The Nyquist loop for this transfer function shows that condition (a) above is satisfied. A similar examination of the Nyquist plot, including the readily computed cable impedance shows that conditions (b) and (c) above are satisfied. Thus the linear analysis indicates stable operation for the system of Fig. 1. This result was confirmed by tests of conditions (a), (b), and (c) individually and by the behavior of the system as a whole, both in the laboratory with a simulated power network for the cable and in the final installation.

One of the most obscure aspects of the power system behavior is that of equilibrium conditions after one or a series of large earth potential

* The analysis briefly summarized here was made by C. A. Desoer.

† In this discussion of simpler systems a power supply consists of only the elements shown in Fig. 2.

‡ In the course of these time-domain measurements, it was quite apparent that the ac control loop could be considered as a linear network only in an approximate sense and thus that the analytical results were primarily useful in interpretation of observed behavior of the system.

disturbances. While the system is stable in the sense that the transient due to a perturbation will disappear in a finite time once the disturbance has been withdrawn, the range of possible equilibrium positions (disregarding the overvoltage protective feature) is extremely wide — from perhaps 1,300 to 2,500 volts at the cable terminals. The upper limit of 2,500 volts is set by the maximum output available from one power plant; this also sets the lower limit of the associate power plant at the far end of the cable. This situation is illustrated diagrammatically in Fig. 9.

The behaviors of the servomechanisms at the two ends of a given cable are nearly enough alike that the repeated introduction of simulated earth potential in the laboratory was found not to disturb substantially the equilibrium point. This was true for earth potentials of either polarity up to 1,000 volts and with these potentials introduced at any point along the artificial cable. A rate of change of earth potential of 20 volts per second was adopted in these tests with the thought that such values would be realistic.

With regard to the long-term stability of the equilibrium condition described above, it is, of course, important that the controls which establish the cable current at the two ends of the same cable be adjusted for very nearly the same value. Unless this is done, the cable voltage at one end will gradually increase or decrease and the voltage at the other end will move equally in the opposite direction. This would eventually

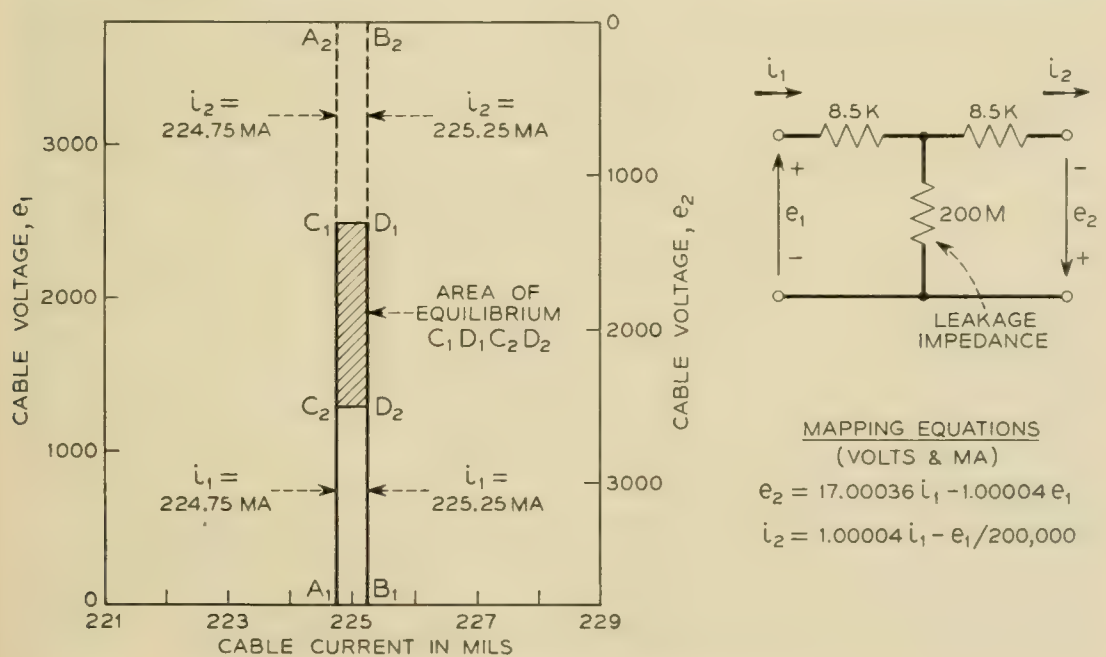


Fig. 9 — Equilibrium diagram.

bring in alarms and necessitate manual readjustment. In this connection, the voltmeters which indicate the drop through the series regulating tubes provide a very convenient magnification of any drift in cable current. The multiplying factor is the effective dc impedance of the regulated system, that is, more than 100,000 ohms. Thus 25 volts, which is an appreciable fraction of the nominal 300 volts across the series regulating tubes, is equivalent to less than 0.25 ma., which is of the order of 0.1 per cent of normal cable current. As a matter of fact, the behavior of this voltage provides the final criterion for precise adjustment of cable current to assure long-term stability.

EQUIPMENT DESIGN

Description

Fig. 10 shows the complete dc equipment for supplying one polarity of power to one cable. Similar equipment provides the opposite polarity to the other cable. The two equipments are located facing each other across a common aisle with their common control bays directly opposite. Regulator 1 on the right is normally operated in parallel with Regulator 2. Regulator 3 on the left is the spare regulator, normally off. The common bay, between Regulators 1 and 2, includes the cable

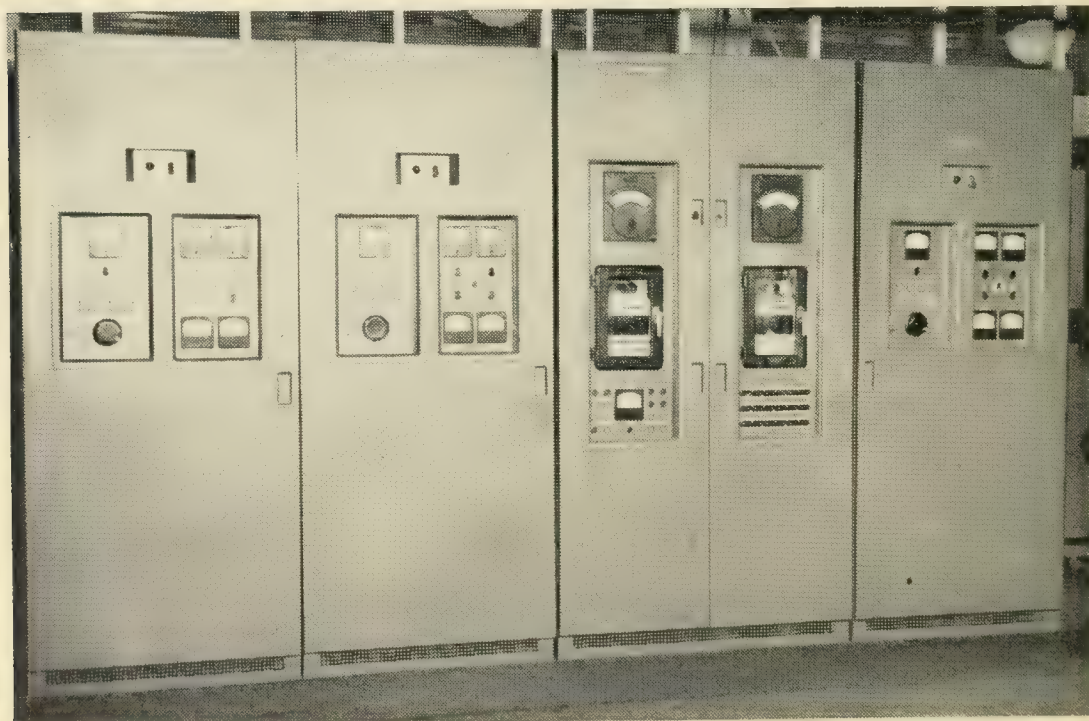


Fig. 10 — DC regulator and common control bays.

termination and power separation filter. This equipment and all the live parts of the circuit, back to the common point to which the switches in the individual regulator bays are connected, is enclosed in a high voltage compartment.

The paralleling control switches are mounted in high voltage compartments in their respective regulating bays and must remain completely enclosed, as their common cable connections are alive during cable operation. These switches have an interrupting capacity of 1 ampere at 3,000 volts, thus providing a large safety factor over the 0.245 ampere maximum load current.

Fig. 11 illustrates some of the special design features built into the equipment to facilitate maintenance. The high voltage compartment shown open at the top is locked whenever the cable is in operation and this protection feature will be described below. Pull-out drawers at the bottom contain metering shunts, a test unit for adjusting voltage and current protection limits, a voltage protection unit, a current protection unit, and an alarm unit. While only one of these compartments is to be pulled out at a time, they are arranged so as not to endanger personnel or to affect service during adjustment when open. Doors are provided on all bays to prevent accidental disturbance of adjustments and to protect against damage to controls.

Corona

The high voltage ac elements of the complete regulator bays were tested for corona with 4,000 volts ac applied, and furthermore, if corona was observed on increasing the applied rms voltage to 5,000 volts, it was required to extinguish when the voltage was reduced to 3500 volts. The maximum acceptable leakage was 20 microamperes at 4000 volts across the circuit (200 megohms). A dc corona requirement of 4000 volts was applied to the dc elements of the regulator bays and 5000 volts for the common bay, with a maximum permissible leakage of 5 microamperes. The higher corona requirements on the common bay were intended to eliminate the necessity for turning down the entire system for repair. A high standard of workmanship is required to provide such performance. There can be no sharp projections and no loose strands of wire. Solder must be applied in such a manner as to obtain a rounded smooth joint and high voltage wiring must be dressed away from exposed grounded metal, bus bars, etc., so that the outer braid (other than polyethylene) does not come in contact with metal.

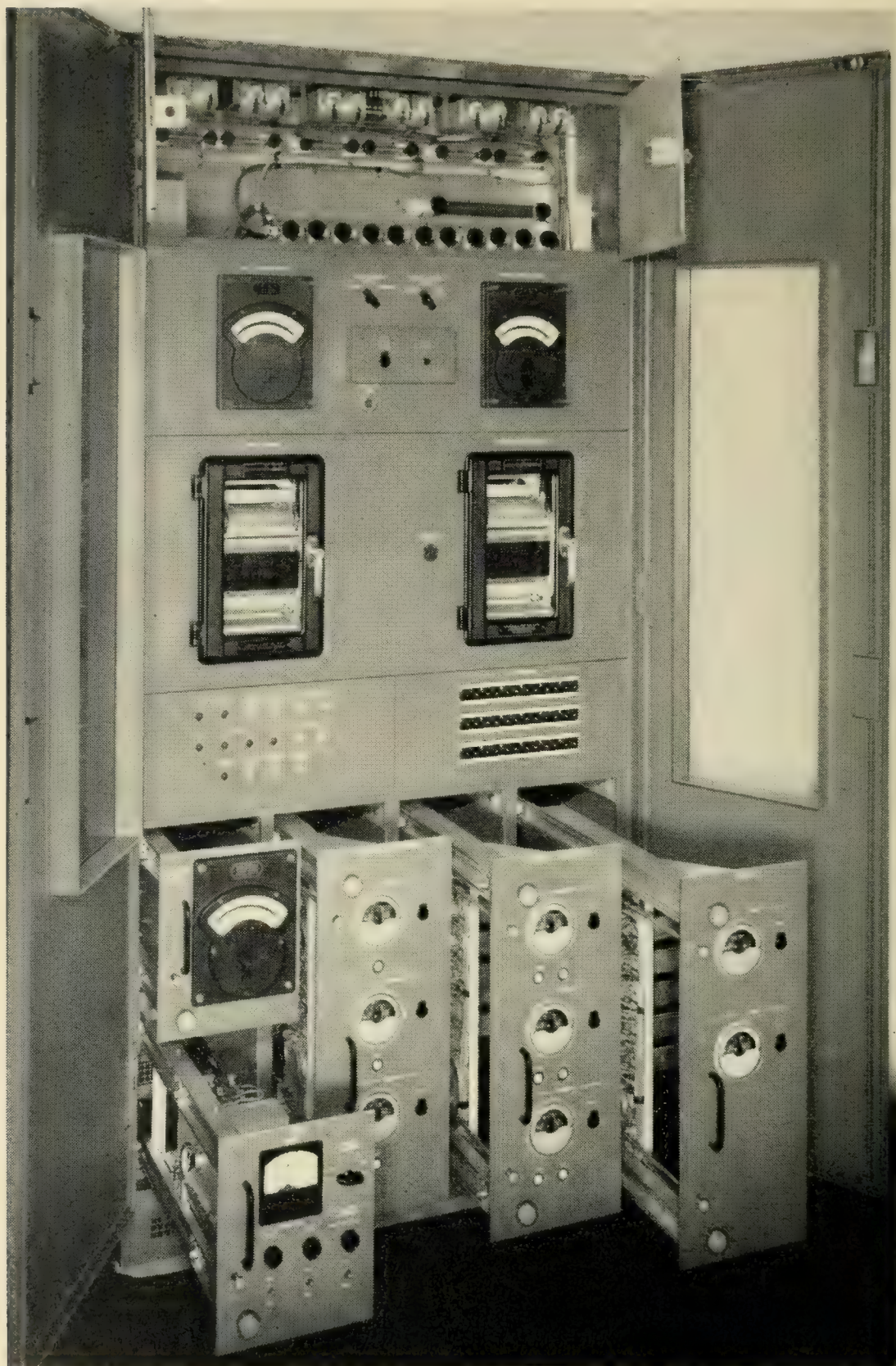


Fig. 11 — Common control bay.

Crosstalk and Outside Interference

In order to meet the severe crosstalk requirements between receiving and transmitting circuits and to guard against feeding office noise potentials into the carrier transmission system, arrangements were made so that office grounds are carefully separated from the outer conductor of the cable and from all circuit elements within the power separation filter. Pickup of external radio-frequency fields by the power separation filters was greatly reduced by completely enclosing in a copper shield the cable terminal and the power separation filter elements nearest to the terminal. The shielding itself and the cans of PSF capacitors and oil-filled coils are connected to the return tape of the cable which is insulated from office ground until it reaches sea water, thus reducing the coupling to the other cable as compared to tying both tapes together at the office or bay frame ground.

Protection of Personnel

A key locking system is provided to safeguard against any hazard to personnel from high voltages. In the common bay, the high voltage compartment can be entered only by operating a switch which shorts the cable to ground and releases a key for the compartment doors. In each regulating bay, the key system assures that the bay is disconnected from the cable and hence from the paralleling power supplies. Where access is required to the interior of any compartment, the key system insures that the ac power to the bay also be switched off.

The test compartment contains pin jacks, provided for maintenance operations which are always performed with the regulator bay connected to a low resistance load. Access to this compartment can be obtained with ac power connected to the bay. However, for such access, the key system enforces the operation of the output disconnect switch, which also transfers the bay to a low-resistance load. Moreover, a mechanical interlock with the autotransformer assures that the test voltages are reduced to safe values.

In addition to its function in protecting personnel, the key system also insures that no more than one regulating bay is disconnected at one time so that continuity of service is protected at all times by two parallel regulators.

FACTORY AND SHIP CABLE POWER

In addition to the above cable power supplies at the ocean terminals, similar dc cable current regulating equipment was designed for use at

the cable factory and abroad the cable ship *Monarch*. Well protected and closely regulated reliable power was considered essential during the cable loading and laying operations. It was necessary to have power on the cable continuously, except when splices were made, in order to detect a fault immediately, to measure transmission characteristics for equalization purposes and finally to alleviate the strain on the glassware and tungsten filaments of the repeater tubes during the difficult laying period.³

REFERENCES

1. H. A. Lewis, R. S. Tucker, G. H. Lovell and J. M. Fraser, System Design for the North Atlantic Link. See page 29 of this issue.
2. T. F. Gleichmann, A. H. Lince, M. C. Wooley and F. J. Braga, Repeater Design for the North Atlantic Link. See page 69 of this issue.
3. J. S. Jack, Capt. W. H. Leech and H. A. Lewis, Route Selection and Cable Laying for the Transatlantic Cable System. See page 293 of this issue.

Electron Tubes for the Transatlantic Cable System

By J. O. McNALLY,* G. H. METSON,† E. A. VEAZIE* and
M. F. HOLMES†

(Manuscript received October 10, 1956)

Electron tubes for use in repeatered underwater telephone cable systems must be capable of operating for many years with a reasonable probability of proper functioning. In the new transatlantic telephone cable system the section of the cable between Nova Scotia and Newfoundland contains repeaters developed by the British Post Office Research Station at Dollis Hill. These repeaters are built around the type 6P12 tube developed at that research station. The repeaters contained in the section of the cable system between Newfoundland and Scotland are of Bell System design and depend on the 175HQ tube developed at Bell Telephone Laboratories.

In this paper the philosophy of repeater and tube design is discussed, and the fundamental reasons for arriving at quite different tube designs are pointed out. Some of the tube development problems and the features introduced to eliminate potential difficulties are described. Electrical characteristics for the two types are presented and life test data are given. Fabrication and selection problems are outlined and reliability prospects are discussed.

INTRODUCTION

Electron tubes suitable for use in long submarine telephone cables must meet performance requirements that are quite different from those imposed by other communication systems. In the home entertainment field, for example, an average tube life of a few thousand hours is generally satisfactory. In the field of conventional land-based telephone equipment, where the replacement of a tube may require that a maintenance man travel several miles, an average life of a few years is considered reasonable. In deep-water telephone cables such as the new transatlantic system, the lifting of a cable to replace a defective repeater may cost several hundred thousand dollars and disrupt service for an extended

* Bell Telephone Laboratories. † British Post Office.

period of time. These factors suggest as an objective for submerged repeaters that the tubes should not be responsible for a system failure for many years, possibly twenty, after the laying of the cable. Such very long life requirements make necessary special design features, care in the selection and processing of materials that are used in the tubes, unusual procedures in fabrication, detailed testing and long aging of the tubes, and the application of unique methods in the final selection of individual tubes for use in the submerged repeaters.

As indicated in the foreword and discussed at length in companion papers, the British Post Office developed the section of the cable system between Clarenville, Newfoundland, and Sidney Mines, Nova Scotia. This part of the transatlantic system uses the 6P12 tube which was developed at the General Post Office (G.P.O.) Dollis Hill Research Station. The submerged portion of this system contains 84 tubes in 14 repeaters. Bell Telephone Laboratories developed the part of the system between Clarenville, Newfoundland, and Oban, Scotland. This section requires 102 repeaters, including 306 tubes, of a type known as the 175HQ. Although a common objective in the development of each of the two sections has been to obtain very long life, the tube designs are quite different.

The Bell System decided on the use of a repeater housing that could be treated as an integral part of the cable to facilitate laying in deep water. The housing is little larger than the cable and is sufficiently flexible to be passed over and around the necessary sheaves and drums. In such a housing the space for repeater components is necessarily restricted. This space restriction, combined with the general philosophy that the number of components should be held to an absolute minimum and that each component should be designed to have the simplest possible structural features, has resulted in the choice of a three-stage, three-tube repeater. In this design, each tube carries the entire responsibility for the continuity of service.

The Post Office Research Laboratories, prior to the development of the transatlantic cable system, had concentrated their efforts on shorter systems for shallow water. The placing of the repeaters on the bottom did not present the serious problems of deep-sea laying, so more liberal dimensions could be allowed for the repeater circuit. A three-stage amplifier was developed which consisted of two strings of three tubes each, parallel connected, with common feedback. The circuit was so designed that almost any kind of tube failure in one side of the amplifier caused very little degradation of circuit performance. This philosophy of having

the continuity of service depend on two essentially independent strings of tubes has been carried over to the repeater design for the Clarendville-Sidney Mines section of the transatlantic cable.

In the Post Office system containing 84 tubes in the submerberd repeaters, five tube failures randomly occurring in the system will result in slightly over fifty per cent probability of a system failure; one tube failure in the 306 tubes in the Newfoundland-Scotland section of the system will result in certain system failure. It is not surprising, therefore, to find the tube designed for the Newfoundland-Scotland section of the cable to have extremely liberal spacing between tube elements in order to minimize the hazards of electrical shorts. This results in a lower transconductance than is found in the tubes designed for the Nova Scotia-Newfoundland link. Other factors in the design will be recognized as reflecting the different operating hazards involved.

Early models of the British Post Office and Bell Laboratories tubes, together with the final tubes used in the cable system, are shown in Fig. 1.

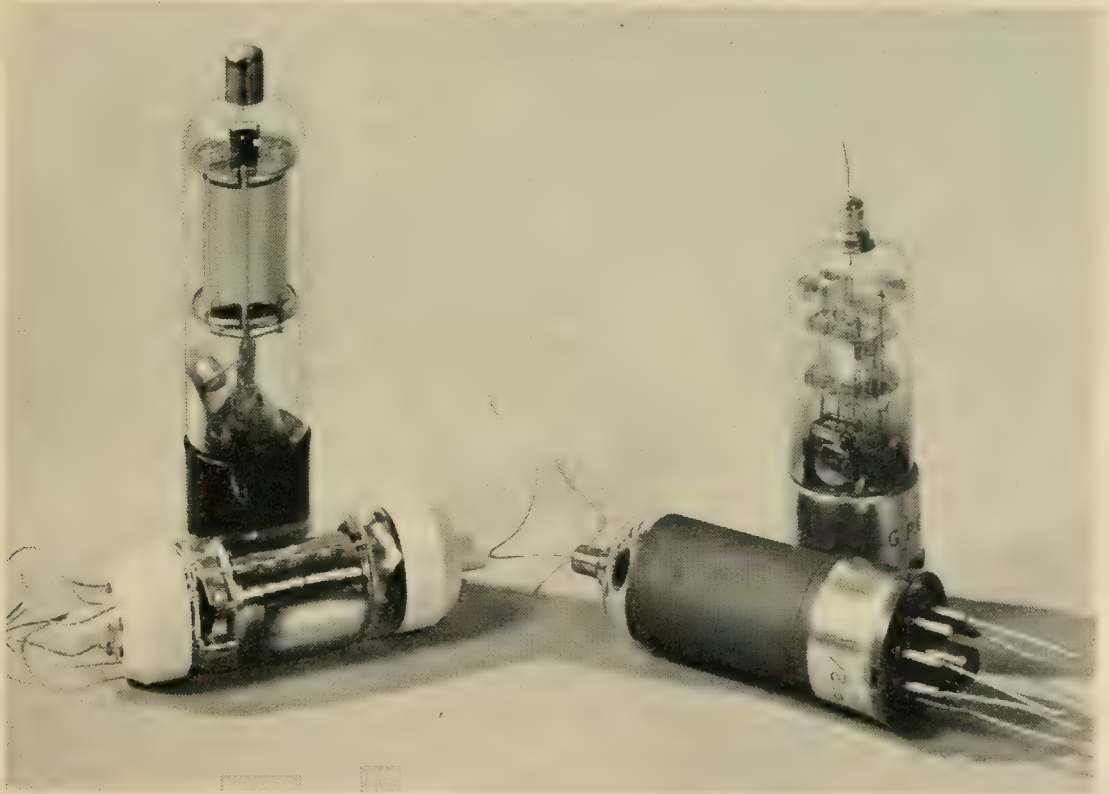


FIG. 1 — The final designs of tubes for the Nova Scotia-Newfoundland section of the cable (right) and for the Newfoundland-Scotland section (left). Early models of each type stand behind the final models.

TUBES FOR THE NEWFOUNDLAND-SCOTLAND CABLE

Early Development Considerations

In Bell Telephone Laboratories, work on tubes for use in a proposed transatlantic cable was started in 1933. This was preceded by a study of what type of tube would best fit the needs of the various proposed amplifier systems and by consideration of what might be expected to give the best life performance.

At the time this project was started, reasonably good tube life had been established for the filamentary types used in Bell System repeaters. Some groups of tubes had average lives of 50,000 or 60,000 hours (6 or 7 years). Equipotential cathode tubes were not then used extensively in the plant, and there was no long life experience with them. However, there appeared to be no basic reason why inherently shorter thermionic life should be expected using the equipotential cathode and there were several advantages in its use. One was the greater freedom in circuit design afforded by the separation of the cathode from the heater. Also there was the possibility of operating the heaters in series and using the voltage drop across the heaters for the other circuit voltages. It was felt, in addition, that the overall mechanical reliability would be greater if the cathode were stiff and rigidly supported.

The first equipotential tubes made were triodes. They were designed for use in push-pull amplifiers wherein continuity of service might be retained in case of a tube failure. This circuit was abandoned in favor of a three tube, feedback amplifier that was the forerunner of the present repeater. The pentode was favored over the triode for this amplifier for obvious reasons, and in 1936 the triode development was discontinued.

Early in the development of the tube three basic assumptions were made. These were, (a) that operation at the lowest practical cathode temperature would result in the longest thermionic life, (b) that operating plate and screen voltages should be kept low, and (c) that the cathode current density should be kept as low as practicable.

The first assumption, concerning the cathode temperature, was based on the observation of life tests on other types of tubes. While the data at the time of the decision were not conclusive, there was definite indication that too high a cathode temperature shortened thermionic life. Little was known about life performance in the temperature range below the values conventionally used.

The second assumption, concerning low screen and plate voltages, had not been supported by any experimental work available at the time of decision. Sixty volts was originally considered for the output stage; this

value was later lowered when other operating conditions were changed. Subsequent results showed that in this range the voltage effects on thermionic life were relatively negligible.

The third assumption, that low cathode current density favored longer thermionic life, affected the tube design by suggesting the use of a large coated cathode area. This implied the use of relatively high cathode power. It was decided early in the planning of the repeater that the voltage drop across the three heaters operated in series would be used to supply part or all of the operating plate and screen potentials. For a 60-volt plate and screen supply, the heater voltage could be as high as 20 volts. A quarter of an ampere was considered a reasonable cable current consistent with voltage limitations at the cable terminals. Thus 5.0 watts were available for each cathode. With this power, a coated area of 2.7 square centimeters was provided. The value of the cathode current, the cathode area, and the interelectrode spacings define the transconductance. Very liberal interelectrode spacings were provided consistent with reasonable tube performance. The original design called for a spacing of 0.040 inch between control grid and cathode. This value was later reduced to 0.024 inch, and a satisfactory design was produced which gave 1,000 micromhos or one milliamperere per volt at a cathode current drain of approximately 2.0 milliamperes. The resulting current density of approximately 0.7 milliamperere per square centimeter is in sharp contrast with values such as 50 milliamperes per square centimeter used currently in tubes designed for the more conventional communication uses. Subsequent data, discussed later, indicate that for current densities of a few milliamperes per square centimeter, the exact value is not critical in its effect on thermionic life.

Subsequent Production Programs

The development of the tube was pursued on an active basis through the years leading up to World War II. During the war development activity essentially stopped. It was only possible to keep the life tests in operation. After the war the development of the tube was completed and a small production line was set up in Bell Telephone Laboratories under the direct supervision of the tube development engineers to make and select tubes for a cable between Key West, Florida, and Havana, Cuba.

This cable turned out to be a "field trial" for the transatlantic cable which was to come later. A total of 6 submerged repeaters containing 18 tubes were laid and the cable was put in operation in June, 1950. The cable has been in operation since this date without tube failure, and

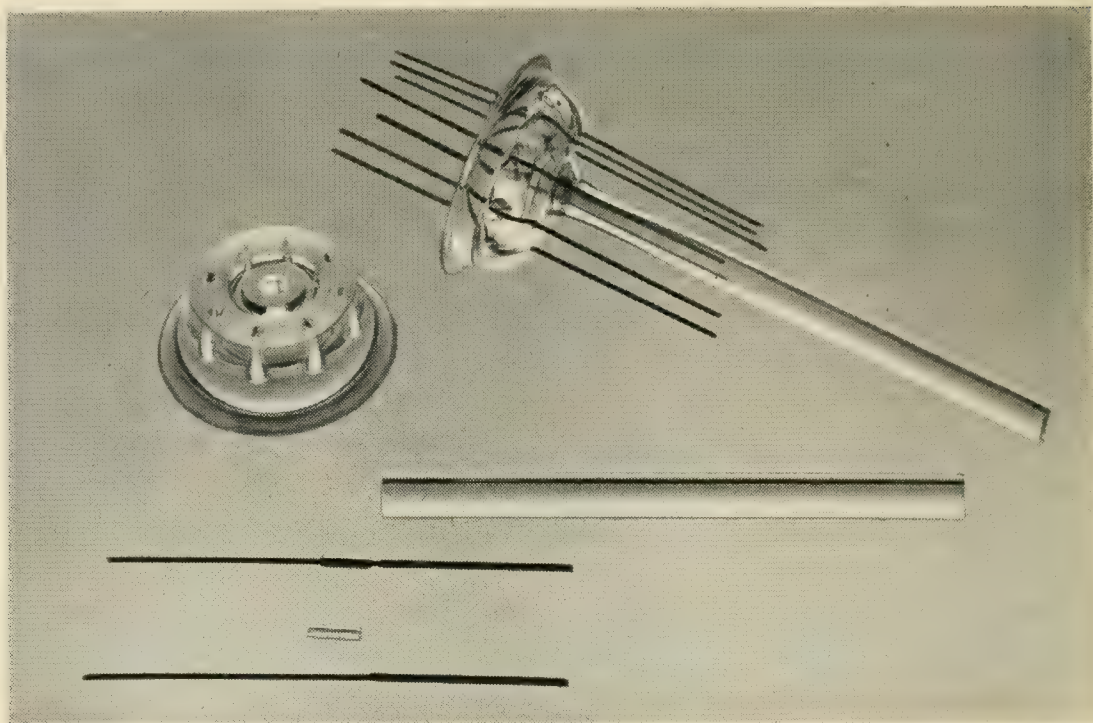


FIG. 2—Parts used in the stem and a finished stem of the 175HQ tube. The separate beading of the leads may be noted.

periodic observations of repeater performance indicate no statistically significant change in tube performance over the 6 years of operation.

Sufficient tubes were made at the same time as the Key West-Havana run to provide the necessary tubes for a future transatlantic cable. These tubes were never used principally because the tubes had been assembled with tin plated leads. Tin plating, subsequent to the laying of the Key West-Havana cable, was found to be capable of growing “whiskers”.¹

In 1953 another production setup was made, also in Bell Telephone Laboratories, for the fabrication of tubes for the Newfoundland-Scotland section of the transatlantic cable. On the completion of this job fabrication was continued to provide tubes for an Alaskan cable between Port Angeles, Washington, and Ketchikan, Alaska. After a pause of several months another run was made to provide tubes for a cable to be laid between California and the Hawaiian Islands.

Mechanical Features

The tube, shown on the left in Fig. 1, is supported in the repeater housing by two soft rubber bushings into which the projections of the two ceramic end caps fit. All leads are flexible and made of stranded beryllium copper which has been gold plated before braiding. Both for

convenience in wiring in the circuit and to hold down the control-grid to anode capacitance, the grid lead has been brought through the opposite end of the tube from the other leads.

A number of somewhat unusual constructional features appear in the tube. The stem on which the internal structure is supported consists of a molded glass dish into which seven two-piece beaded dumet leads are sealed. The parts used in a stem, and also a finished stem, are shown in Fig. 2. It is usual to embed the weld or "knot" between the dumet and nickel portions of the lead in the glass seal to provide more structural stiffness. This has not been done in this stem because it was believed that a fracture of a lead at the weld could be detected more easily if it were not supported by the seal. It might be questioned why the modern alloys and flat stem structure have not been used. It is to be remembered that one gas leak along a stem lead would disable the system, and experience built up with the older materials provides greater assurance of satisfactory seals.

The structure of the heater and cathode assembly is unique, as may be seen in Fig. 3. A heater insulator of aluminum oxide is extruded with 7 holes arranged as shown. This insulator is supported by a 0.025 inch molybdenum rod inserted in the center hole. The heater consisting of about 36 inches of 0.003 inch tungsten is wound into a helix having an outside diameter of 0.013 inch. After dip coating by well known techniques the heater is threaded through the 6 outer holes in the insulator. A suspension of aluminum oxide is then injected into the holes in the insulator so that on final firing the heater becomes completely embedded.

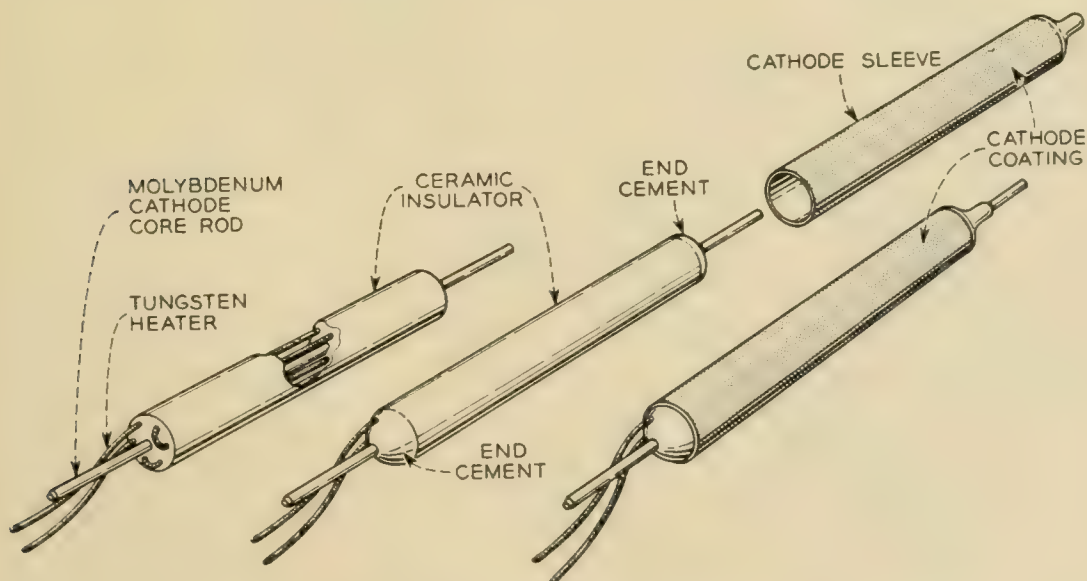


FIG. 3 — Heater, heater insulator and cathode assembly of the 175HQ tube.

The cathode sleeve, which is necked down at one end as shown in Fig. 3, is slipped over the heater assembly and welded to the central molybdenum rod which becomes the cathode lead. By this means a uniform temperature from end to end of the cathode is obtained. Under normal operating conditions the heater temperature is approximately 1100°C , which is very considerably under the temperature found in other tubes.

Connection of the heater to the leads from the stem presented a serious design problem. Crystallization of tungsten during and after welding and mechanical strains developed by thermal expansion frequently are the causes of heater breakage. This problem was successfully overcome by the means illustrated in Fig. 4. Short sections of nickel tubing are slipped over the cleaned ends of the heater coil and matching pieces of nickel wire are inserted as cores. These parts are held together by tack welds at the midpoints of the tubing. The heater stem leads are bent, flattened and formed to receive the ends of the heater, which are then fastened by welds as indicated in the drawing.

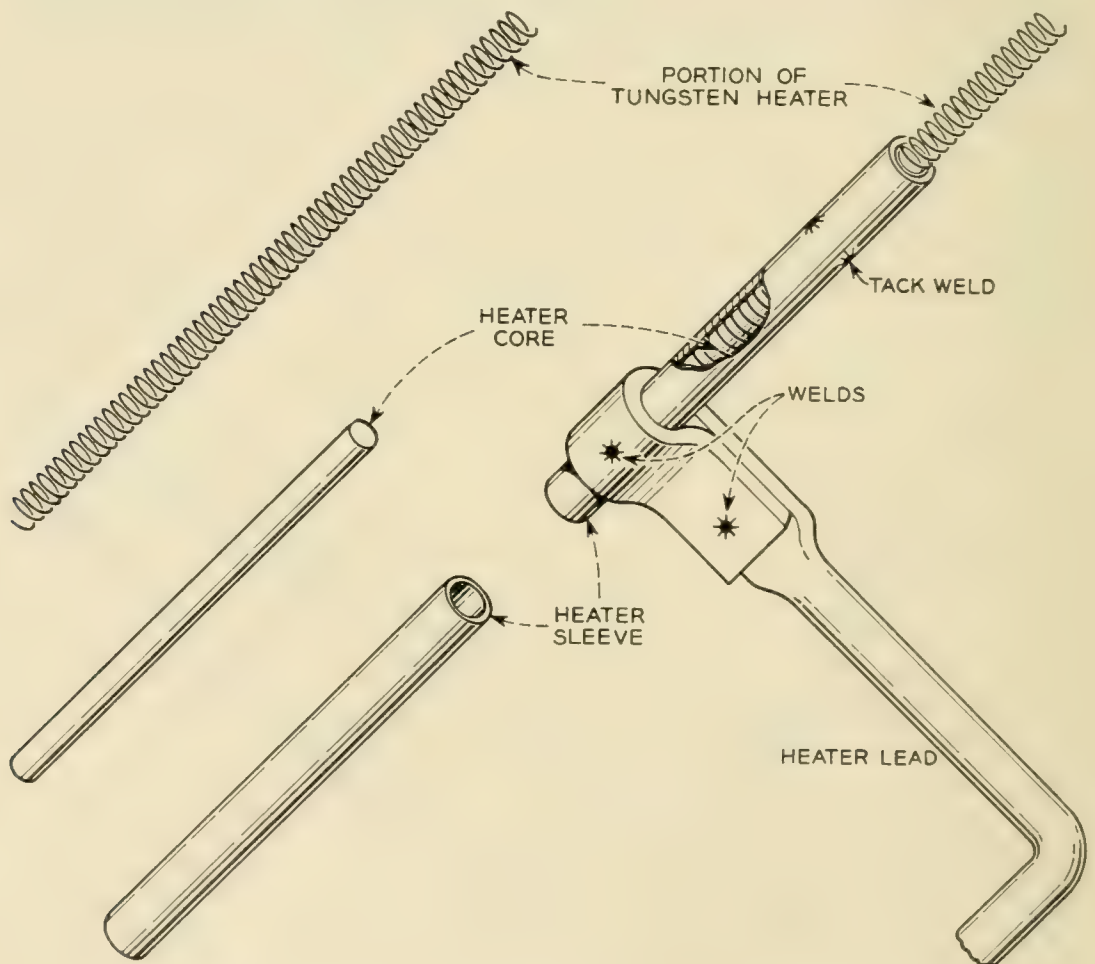


FIG. 4 — Heater tabbing arrangement of the 175HQ.

A serious attempt has been made in the design of the tube to hold the number of fastenings that depend entirely on one weld to an absolute minimum. The grids are of conventional form in which the lateral wires are swaged into notches cut in the side rods or support wires. The side rods as well as the lateral wire are molybdenum. This produces grids which are considerably stronger than those using more conventional materials.

The upper mica is designed to contact the bulb and the bulb is sized to accurate dimensions to receive and hold the mica firmly. The tube in its mounting will withstand a single 500*g* one millisecond shock without apparent changes in mechanical structure or electrical characteristics. It is estimated from preliminary laying tests that accidental or unusual handling would rarely result in shocks exceeding 100*g*.

Electrical Characteristics and Life

The average operating electrical characteristics for the 175HQ tube are given in Table I, and a family of plate-voltage versus plate-current curves for a typical tube is given in Fig. 5 for a region approximating the operating conditions.

The development of a long-life tube offers good opportunities to observe effects which are more likely to be missed where shorter lives are satisfactory. For example, some of the earliest tubes made, after 20,000 hours on the life racks, began to show a metallic deposit on the bulbs.

TABLE I — AVERAGE OPERATING ELECTRICAL CHARACTERISTICS FOR THE 175HQ TUBE

	Stages 1 & 2	Stage 3
Heater Current.....	220	217 milliamperes
Heater Voltage.....	18.2	18.4 volts
Heater Power.....	4.0	4.0 watts
Control-Grid Bias.....	−1.3	−1.4 volts
Screen Voltage.....	38	40 volts
Plate Voltage.....	32	51 volts
Screen Current.....	0.3	0.3 milliamperes
Plate Current.....	1.3	1.4 milliamperes
Transconductance.....	980	1010 micromhos
All Stages		
Capacitances (cold, with shield)		
Input Capacitance.....		9.2 μμf
Output Capacitance.....		15.6 μμf
Plate to Control-Grid Capacitance.....		0.03 μμf

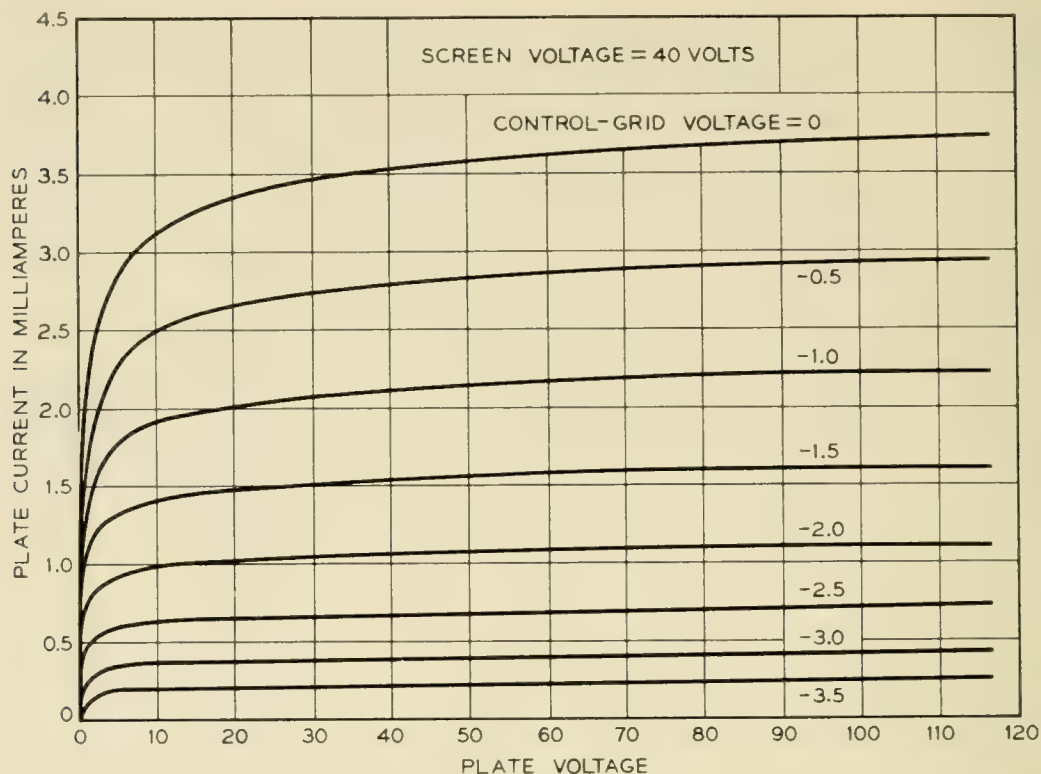


FIG. 5 — Typical plate voltage-plate current characteristics for a type 175HQ tube.

Immediate concern for the lowering of insulation resistance across mica spacers prompted an investigation. The source was traced to the use of plates made from a grade of nickel from which magnesium as a contaminant was evaporating. A change was made to molybdenum which has been used successfully since that experience.

The effect on the thermionic life of operating at different cathode current densities was of interest. Life tests were started in which the cathode current drain in one group of tubes was approximately 7.5 milliamperes (2.8 ma/cm^2) and in another group the average cathode current was 0.6 milliamperes (0.2 ma/cm^2). The results presented in Fig. 6 after 120,000 hours, or approximately 14 years, show that at the cathode temperature of approximately 710°C^* selected for the test, there is practically no current density effect in this 12 to 1 current range. The circles indicate average values, while the dots and crosses at each test point show the positions of the extreme tubes of the group.

Similar life tests set up to show the effect of operating at plate and screen voltages of 60 volts as compared to 40 volts indicate no essential differences in performance after 8 years of operation.

* All cathode temperatures referred to in this paper are "true" temperatures, not uncorrected pyrometer temperatures.

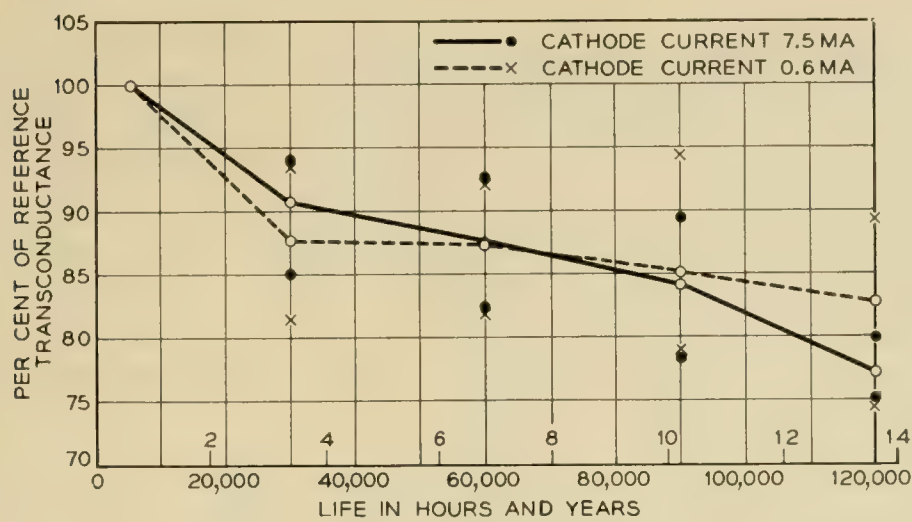


FIG. 6 — Results of life tests on eighteen 175HQ tubes operating at two different current densities.

The cathode temperature is one of the most critical operating variables affecting thermionic life. As mentioned above, the early development objective was a cathode power of 5.0 watts which corresponded to 710°C for the cathode design used at that time. The results of operating at this condition are illustrated in Fig. 7. No tubes have been lost from the test where the direct cause has been failure of emission. Several tubes were lost because of mechanical failure resulting from design defects which were subsequently corrected. It will be observed

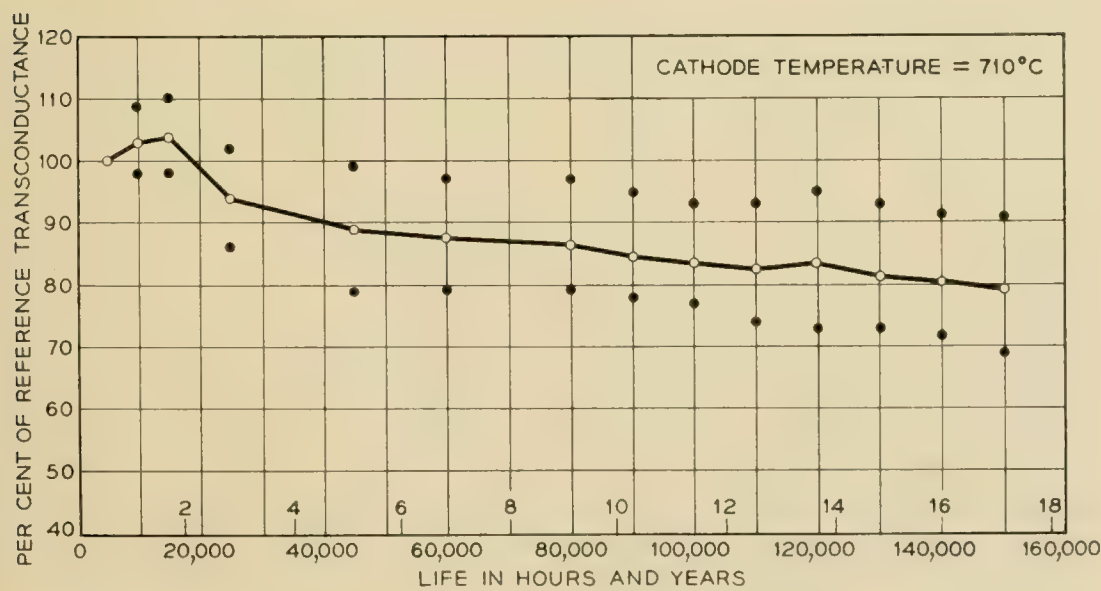


FIG. 7 — Results of life tests on sixteen 175HQ tubes operating at a cathode temperature of 710°C.

that at the end of 17 years the average transconductance is 80 per cent of its original value, and the poorest tube has dropped to 69 per cent. There is reason to believe that test set difficulties may very well account for a large part of the variation shown in the first three years.

The cathode coatings used in all experimental and final tubes for the Newfoundland-Scotland link of the transatlantic cable are the conventional double carbonate coatings. The cathode base material is an International Nickel "220" nickel. The particular melt used for the transatlantic cable is known as melt 84. A typical analysis for melt 84 nickel cathodes is given in Table II.

TABLE II — TYPICAL ANALYSIS OF INCO 220 NICKEL CATHODE MELT 84
(Analysis made prior to hydrogen firing)

Impurity	Per Cent	Impurity	Per cent
Aluminum	0.008	Manganese	0.11
Boron	<0.004	Silicon	0.033
Cobalt	0.46	Titanium	0.032
Chromium	<0.005	Oxygen	0.0001
Copper	0.028	Sulphur	0.0016
Iron	0.093	Carbon	0.058
Magnesium	0.046		

The relatively high carbon content (0.058 per cent) of melt 84 cathode nickel is capable of producing excessive reduction of barium in the cathode coating.^{2, 3} A treatment in wet hydrogen, prior to coating, at 925°C for 15 minutes reduces the carbon in the cathode sleeve to about 0.013 per cent.

Melt 84 was as close as was obtainable in composition to melts 60 and 63 previously used for the Key West-Havana tubes. The results of up to five years of life testing were thus available on materials of very similar composition.

One common cause of tube deterioration with life is the result of formation of an interface layer on the surface of the cathode sleeve. It is known that the rate of development of this layer depends in a complex way on the chemical composition of the nickel cathode core material. The effect of such a layer is to introduce a resistance in series with the cathode. This results in negative feedback and reduces the effective transconductance. Since the effect of a given feedback resistance in this location is proportional to transconductance, the relatively low value for the 175HQ tube tends to minimize this feedback effect. In addition, the low cathode temperature tends to reduce the rate of formation of interface resistance, and the relatively large cathode area tends to further minimize the effects. The final decision to use melt 84 was based on ac-

celerated aging tests which showed it to be superior to melts 60 and 63 from an interface standpoint.

The interface problem will be discussed further in a later section.

As the development of the tube proceeded, both the processing of the parts and the cleanliness of the mount assembly were improved and the cathode emission level increased. Life tests indicated that better thermionic life might be obtained by operating at a lower cathode temperature. Accordingly a cathode power of approximately 4.0 watts was adopted, which corresponds to a temperature of 670°C . A life test, now 45,000 hours or about 5 years old, shows the results in Fig. 8 of operating groups of tubes at three different cathode temperatures. This is a well controlled test in that the tubes for the three groups were picked from tubes having common parts and identical fabrication histories. It may be noted that the average of the 725°C lot has lost approximately 5 per cent of the initial transconductance, whereas the 4.0 watt group after about 5 years has lost essentially none of its transconductance. The 3.0

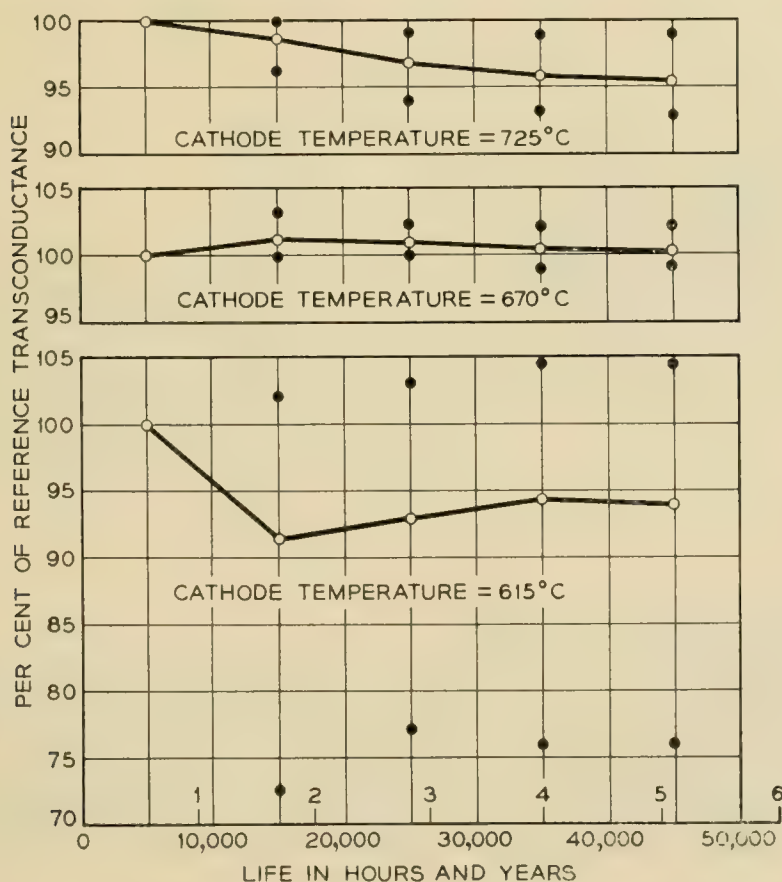


FIG. 8 — Results of operating thirty-six 175HQ tubes divided equally among three different cathode temperature conditions. For each of the curves the cathode core material used was half from melt 60 and half from melt 63. The conditions in cable operation are essentially those represented by the center curve.

watt (615°C) group shows serious instabilities in its performance. In some of the tubes the cathode temperature has not been sufficiently high to provide the required emission levels.

The design of the repeaters in the Newfoundland-Scotland section of the cable is such that reasonably satisfactory cable performance would be experienced if the transconductance in each tube dropped to 65 per cent of its original value. The life test performance data presented in Figs. 7 and 8, and other tests not shown, indicate that operation of the 175HQ tubes in the transatlantic cable at approximately 4.0 watts will assure satisfactory thermionic performance for well over 20 years.

Mention was made that cleanliness in the assembly of the mounts was a factor which affected thermionic activity. Interesting evidence supporting this view was obtained during the fabrication of tubes for the Key West-Havana cable. The quality control type of chart reproduced in Fig. 9 shows the average change in transconductance between two set values of heater current for the first 5 tubes in each group of approxi-

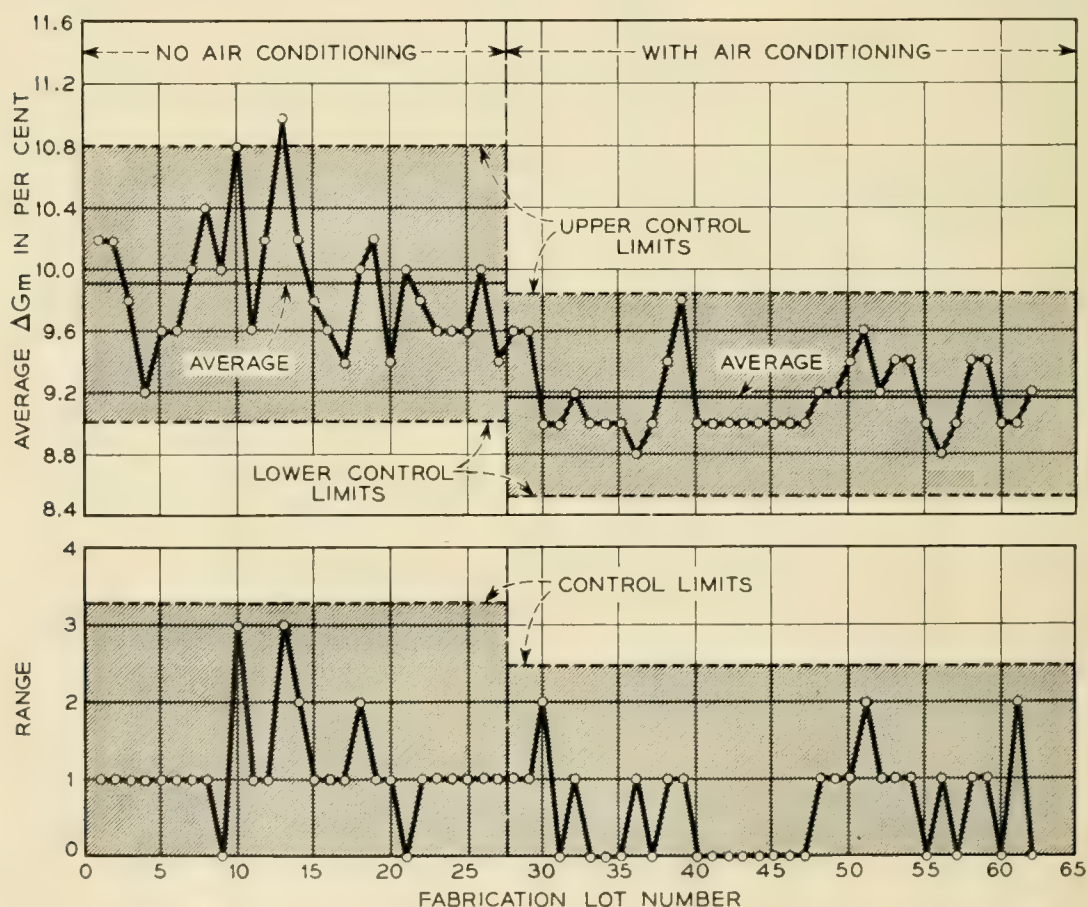


FIG. 9 — Control chart showing the effect of air cleaning on the cathode activity level. The per cent change in transconductance for normal heater current and a value 20 per cent lower is used as the measure of cathode performance.

mately 28 tubes made. The data were taken after 5,000 hours of aging. A sharp improvement in thermionic emission was noted at a point on the chart where about one half of tubes had been fabricated. An examination of the records, which are very carefully maintained, disclosed that the windows of the assembly room were sealed and air cleaning and conditioning was put in effect at the point indicated on the chart. No other changes in processing or materials occurred at this time. A second definite improvement in thermionic emission occurred when the work was moved from the location in New York City to the new and better controlled environment at Murray Hill in New Jersey.

Fabrication and Selection

All assembly operators on the 175HQ tube program wore nylon* smocks to keep down the amount of dust and lint that might otherwise leave their clothing and get into the tubes. Rayon acetate gloves were worn when handling parts as a protection against perspiration. Rubber finger cots did not prove satisfactory because they covered too little area and once contaminated they did not absorb the contaminant.

The tubes for the Newfoundland to Scotland section of the cable were made at Bell Telephone Laboratories under the extremely close engineering supervision of many of the original development engineers. All materials going into the tubes were carefully checked, and wherever possible they were tried out in tubes. Experience under accelerated aging conditions was obtained before these materials were used. For example, although during the development period all glass bulbs were used as received without any failures resulting, less than one-quarter of the bulbs bought for the actual cable job passed the inspection requirements. Each batch of heaters was sampled and results obtained on intermittent and accelerated tests before approval for use.

The fabrication of the tube was carried out with extreme care by operators especially selected for the job. If normal commercial test limits were applied to the tubes after exhaust, the yield from acceptable mounts would have been about 98 to 99 per cent. Yet only approximately one out of every seven tubes pumped was finally approved for cable use. All tubes were given 5,000 hours aging and electrical tests were made at six different times during this period. The results weighed heavily in the final selection. For example, a correlation between thermionic life and gas current had been established during the tube develop-

* Trade name for DuPont polyimide fibre.

ment period, and only tubes having control-grid currents due to gas of less than 5×10^{-11} ampere were acceptable. This corresponds to a gas pressure of approximately 2×10^{-7} mm of Hg. Very thorough mechanical inspections after the 5,000-hour aging were made to insure that there were no observable mechanical deviations that could cause trouble. The history of each group of 28 tubes, from which prospective candidates for the cable were selected, was reviewed to see if any group abnormalities were found. In case a suspected trait was seen, all tubes in the group were ruled out for cable use.

As an aid in the selection of tubes for the cable, all pertinent data were put on IBM cards. It was then possible to manipulate and present the data in many very helpful ways that would have otherwise been wholly impractical from time and manpower considerations. An over-all total of about half a million bits of information was involved.

Reliability Prospects

Questions are frequently raised concerning the probability of tube failures in the system. There are two areas into which failures naturally fall—catastrophic failures and the type of failure caused by cumulative effects such as the decay of thermionic activity, development of primary emission from the control-grid, or the build-up of conductance across mica insulators or glass stems.

The catastrophic failures might include such items as open connections caused by weld failures or fatigue of materials, short circuits caused by parts of two different electrodes coming into contact or being bridged by conducting foreign particles and gas leaks through the glass or along stem leads. Fortunately these failure rates have been lowered to a point where there are no sound statistical data available in spite of the substantial amount of life testing that has been done. In approximately 4,800 tubes made to date, there have been four failures that were not anticipated by the inspections made. All four of these failures were of different types and occurred either at or before 5,000 hours of life. All four were of types more apt to occur during the early hours of aging and handling.

Of the cumulative types of failure, life testing has indicated no apparent problem with either the growth of insulation conductance or primary emission from the grids. As indicated earlier in the paper, thermionic life results are such that there is reason to be optimistic that no failures will occur in 20 years.

TUBES FOR THE NOVA SCOTIA-NEWFOUNDLAND CABLE

*Trend of Tube Development in British Submarine Repeater Systems**Early Use of Commercial Receiving Tubes*

The development of submerged telephone repeaters in Britain has taken a somewhat different course from that followed in the United States. Off North America, deep seas are encountered as soon as the continental shelf has been passed. Consequently emphasis has been placed from the beginning on the design of repeaters for ocean depths. In Britain, separated from many countries by only shallow seas, it was natural for development to start with a repeater specially designed for shallow water. Such a repeater was laid in an Anglo-Irish cable in 1944.

The tubes used in the amplifier of this repeater were normal high transconductance commercial pentodes type SP61. These tubes were known, from life test results, to last at least for two years under conditions of continuous loading. Their performance in the first and subsequent early repeaters exceeded all expectations. So far one tube has failed, and this from envelope fracture, after a period of four years service. There remain 23 of this type on the sea bed with a service life of five or six years, and 3 tubes which have survived ten years.

All these SP61 tubes were part of a single batch made in 1942 and their performance set a high standard. It was, however, found that subsequent batches did not attain the same standard set by the 1942 batch. In 1946, therefore, the British Post Office was faced with the fact that future development of the shallow water system of submerged repeaters was dependent on the production of a tube type which could take the place of the 1942 batch of SP61 tubes. This situation led to the formation of a team at Dollis Hill whose terms of reference were, specifically, to produce the replacement tube and, generally, to study the problems presented by the use of tubes in submerged repeaters. Apart from changes in the specific requirements, these terms of reference have remained unchanged from that day to this.

Replacement by the G.P.O. 6P10 Type

Coincident with the rapid exhausting of stocks of satisfactory SP61 tubes for submerged telephone systems, there arose the need for a tube type for a submerged telegraph repeater. This latter requirement was complicated by the fact that the telegraph cable was subject to severe overall voltage restrictions which precluded the 630 ma heater current required for the 4 watt cathode of the SP61. In order to avoid production

of one tube type for telephone systems and another for telegraph, it was decided that the replacement for the SP61 should have a 2 watt cathode with a 300 ma heater.

During the years 1944 and 1945 a very successful miniature high slope pentode, the CV138, was produced for the armed services. The electrical characteristics of this tube were superior to those of the SP61 and, in addition, it used a 2 watt cathode. It was therefore decided to base the replacement tubes, electrically, on the CV138, whilst, at the same time, retaining freedom to amend the mechanical features in any way which might seem to favor the specific requirements of submerged repeater usage, in particular, maintenance of the level of transconductance unchanged for long periods. Consequently three major mechanical changes were made at the outset of the project. The miniature bulb of the CV138 was replaced by one of normal size (approximately 1 inch diameter and $2\frac{1}{2}$ inches long). This was done to reduce the glass temperature and so reduce gas evolution. At the same time a normal press and drop seal were substituted for the button base and ring seal of the CV138, as it was felt that, with the techniques available, the former would be more reliable than the latter. With the use of a normal press there immediately followed a top cap control grid connection, so producing a double-ended tube in place of the single-ended CV138.

These three modifications and a number of major changes to improve welding and assembly techniques led to the G.P.O. type known as the 6P10 [A pentode (P) with a 6.3-volt heater (6) of design mark 10]. The 6P10 replaced the SP61 in the 18 shallow water repeaters laid in various cables after 1951. There are, therefore, 54 tubes type 6P10 in service on the sea bed with periods of continuous loading ranging from two to four years. There has been one failure due to a fractured cathode tape, and one other repeater was withdrawn from service to investigate a high-frequency oscillation associated with a tube. The oscillation cleared, however, before the cause could be identified.

The first eighteen 6P10 type tubes used in repeaters had conventional nickel cathode cores. Appreciation of the problem of interface resistance led to the use of platinum as a core material for the following 36 tubes. The steps leading to this radical change of technique will be described later.

Development of the 6P12 for Long Haul Systems

Although submerged repeater development started naturally in Britain with shallow-water systems, it was inevitable that attention should

ultimately turn towards trans-ocean cables. The tube requirements for such long haul systems differ from those for short haul shallow-water schemes in that operation at a lower anode voltage is essential. A new tube to replace the 6P10 was therefore unavoidable.

By the time emphasis started to shift in Britain from shallow-water to deep-sea systems some considerable experience had been gained at Dollis Hill on the production techniques required to fit a 6P10 type tube having a platinum cathode core for submerged repeater usage. When it became apparent that a new tube had to be designed for the first long haul system, it was resolved to retain as much as possible of the 6P10 structure in order to take full advantage of familiar techniques. The 6P10 was therefore redesigned for 60-volt operation simply by a major adjustment to the screen grid position and minor alterations elsewhere. The new tube became known as the 6P12 and was used in seven repeaters installed in the Aberdeen-Bergen cable. This scheme was regarded as a proving trial for the Newfoundland-Nova Scotia section of the transatlantic project.

It has always been appreciated that the use of high transconductance tubes using closely-spaced electrodes will involve a higher liability to mechanical failure by internal short circuits. Practical experience in shallow-water schemes, where repeater recovery is a comparatively cheap and simple operation, has shown however that such risks seem to be outweighed by the economic advantage accruing from a tube capable of wider frequency coverage. In point of fact a failure by internal short circuit has not yet materialized on any shallow-water system.

This background of experience explains the British choice of a high transconductance tube for deep-sea systems, but the greater liability to mechanical failure is acknowledged by use of parallel amplifiers. Confidence in this policy has been increased by the successful operation of the Aberdeen-Bergen system.

Problems of Development of the 6P12 Tube

The main preoccupation of the thermionics group at Dollis Hill since 1946 has been a study of the electrical life processes of high transconductance receiving tubes. This effort has led to a conviction that all changes of electrical performance have their origin in chemical or electro-chemical actions occurring in the tube on a micro- or milli-micro scale of magnitude. The form of change of most importance to the repeater engineer is decay of transconductance and this will be considered in brief detail as typical of the development effort put into the 6P12 tube.

Transconductance decay in common tubes results from two separate

and distinct chemical actions occurring in the oxide cathode itself. Both actions are side issues in no way essential to the basic functioning of the cathode and it seems probable that both can be eliminated if sufficient understanding of their nature is available. The first action is the growth of a resistive interface layer between the oxide matrix and its supporting nickel core, discussed briefly in an earlier section. This effect is assumed to be due to silicon contamination of the nickel core metal.



The resistance of the layer of barium orthosilicate rises as it loses its barium activator and approaches the intrinsic state. The effect of the interface resistance is to bring negative feedback to bear on the tube with resulting loss of transconductance. The second deleterious action is loss of electron emission from the oxide cathode by direct destruction of its essential excess barium metal by oxidizing action of residual gases. Such gases result from an imperfect processing technology.

These two problems have been approached in the 6P12 tube in a somewhat novel manner. The conventional nickel core is replaced by platinum of such high purity (99.999 per cent) that the possibility of appreciable interface growth from impurities can be disregarded. The only factor to be considered is the appearance of high resistive products of a possible interaction between platinum and the alkaline earth oxides. Batch tests over a period of 30,000 hours have failed to show any sign whatever of such an action occurring and workers at Dollis Hill now regard the pure platinum-cored tube as free from the interface resistance phenomenon.

The problem of avoiding gas deactivation of the cathode is a more difficult one and so far has been reduced in magnitude rather than eliminated. It is now appreciated that the dangerous condition arises from "gas generators" left in the tube and not from a true form of residual gas pressure left after seal-off from the pump. These gas generators are solid components of the tube which give off a continual stream of gas over a prolonged period of time. The gas evolution rates are usually so small that they cannot be detected by reverse grid current measurement but they tend to integrate gas by absorption on the cathode and to destroy its activity. The gas generators are usually of finite magnitude and, depending mainly on diffusion phenomena, evolve gas at a rate which falls in roughly exponential fashion with time. The probability of transconductance failure is therefore highest in early life and tends to lessen with time as the generators move to exhaustion.

One particularly useful feature of the platinum-cored cathode is its freedom from core oxidation during gas attack and this leaves the tube

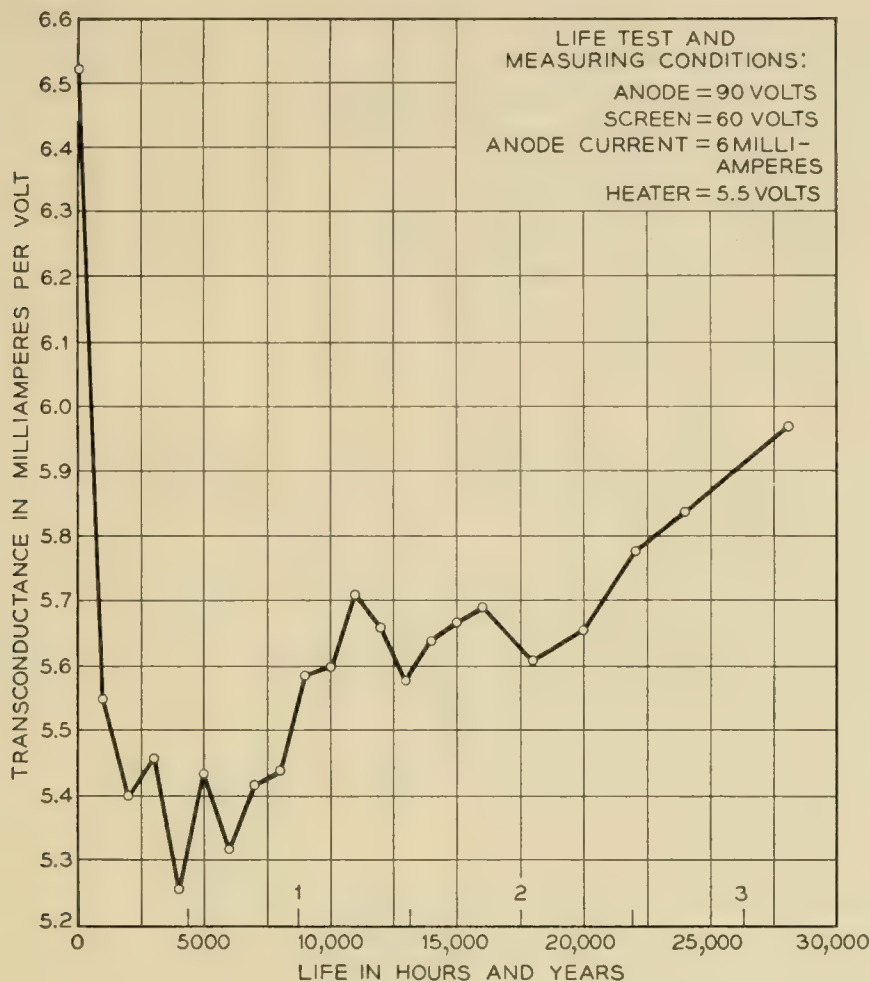


FIG. 10 — Behaviour of a group of 50 tubes deliberately left with a “gas generator” (tube type 6P12).

free to recover from transconductance failure when the gas attack has passed. In Fig. 10 is shown the behaviour of a group of 50 tubes which have been deliberately left in possession of a component capable of generating carbon-monoxide over a prolonged period of time. The curve shows the characteristic recovery of a platinum-cored oxide-cathode with the gradual passing of what is thought to be a typical gas attack.

One problem that has attracted much attention at Dollis Hill is the actual manner in which a platinum-cored cathode recovers from a gas attack. The mechanism must involve the dissociation of a small fraction of the oxide cathode itself with the retention of barium metal in the oxide lattice and the evolution of oxygen. That such an essential mechanism does in fact exist has been proved by the slow accumulation of barium metal in the platinum core. This accumulation takes the form of a distinctive alloy of barium and platinum and only occurs when the cathode

is passing current. The barium regenerative process seems therefore to be electrolytic in nature and, depending only on current flow and a stock of oxide, would appear to be virtually inexhaustible.

These few remarks are perhaps sufficient to give some idea of the lines on which the British research effort has run during the past decade. More detailed descriptions have already been presented elsewhere.^{4, 5, 6, 7}

Electrical and Mechanical Characteristics

Electrical Characteristics

The main electrical characteristics of the 6P12 are shown in Figs. 11 and 12. The heater voltage used for both sets of curves is 5.5 volts, the same value as that used in the British amplifier.

Fig. 11 shows the change of transconductance with anode current, with screen voltage as parameter. An anode voltage of 40 volts and a suppressor voltage of zero correspond with the static operating conditions in the first two stages of both the Aberdeen-Bergen and the British

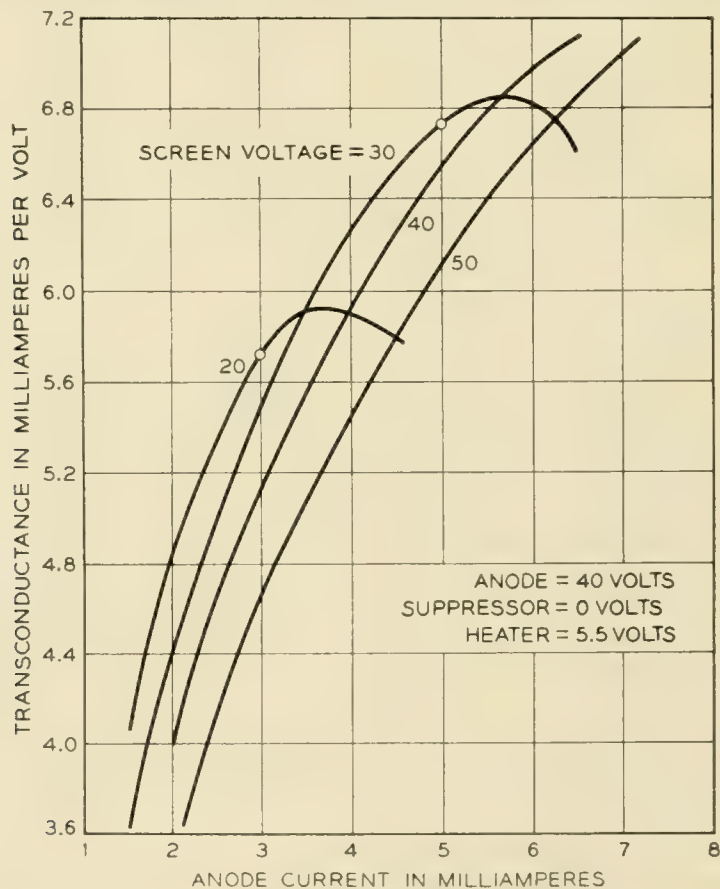


FIG. 11 — Typical transconductance-anode current characteristics for a type 6P12 tube (No. 457/6).

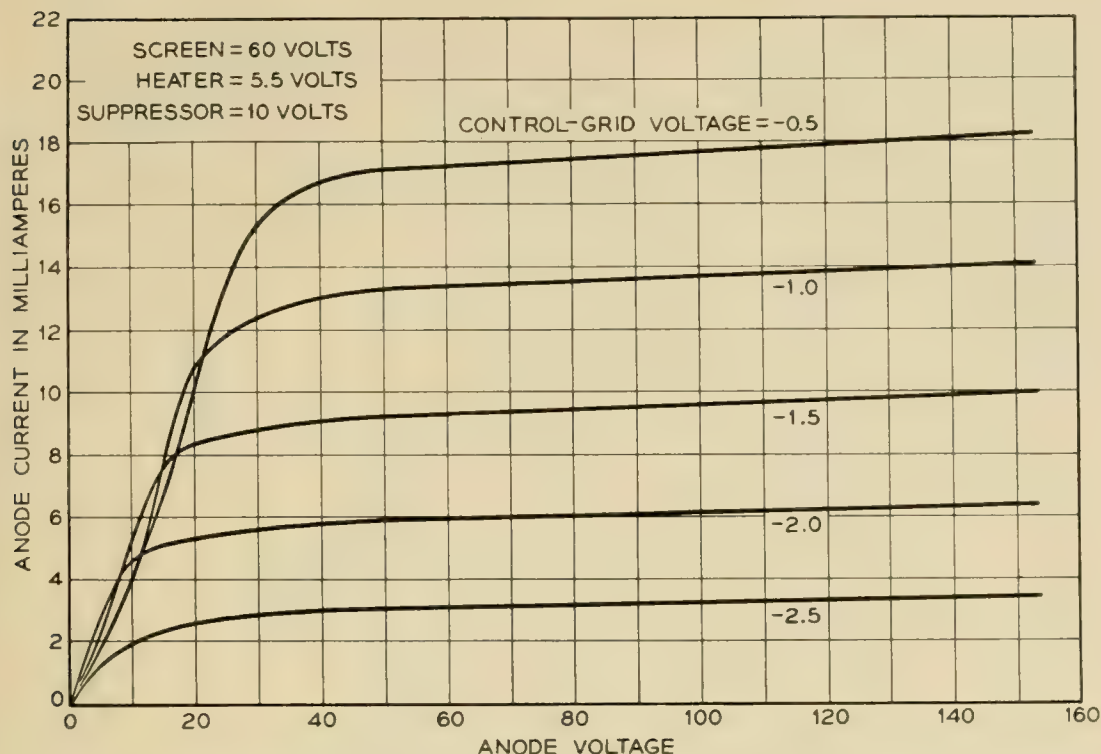


FIG. 12 — Typical anode voltage-anode current characteristics for a type 6P12 tube (No. 457/6).

transatlantic telephone (T.A.T.) amplifiers. The screen voltage and anode current of the first two stages of the British T.A.T. amplifier were chosen to be 40 volts and 3 ma respectively.

Fig. 12 shows the normal anode voltage-anode current characteristics for conditions corresponding to the output stage of the amplifier (static operating point, anode voltage = 90, screen voltage = 60, anode current = 6 ma). A final electrical characteristic worthy of comment is the level of reverse grid current. For all specimens tested at the time of selection, after about 4,000 hours life test, the level is very low, about 100 micromicroamperes per milliampere of anode current.

Life Performance

The life performance of the 6P12 is still a matter for conjecture. The only concrete evidence available is the behaviour of a group of 92 tubes which were placed on life test some three years ago. The change of the average transconductance of this group (with anode current constant at 6 ma) is shown in Fig. 13. It may be clearly appreciated that there is no definite trend over the past year which permits any firm prediction of life expectancy. Examination of other tube characteristics is equally unproductive from the point of view of prediction of failure.

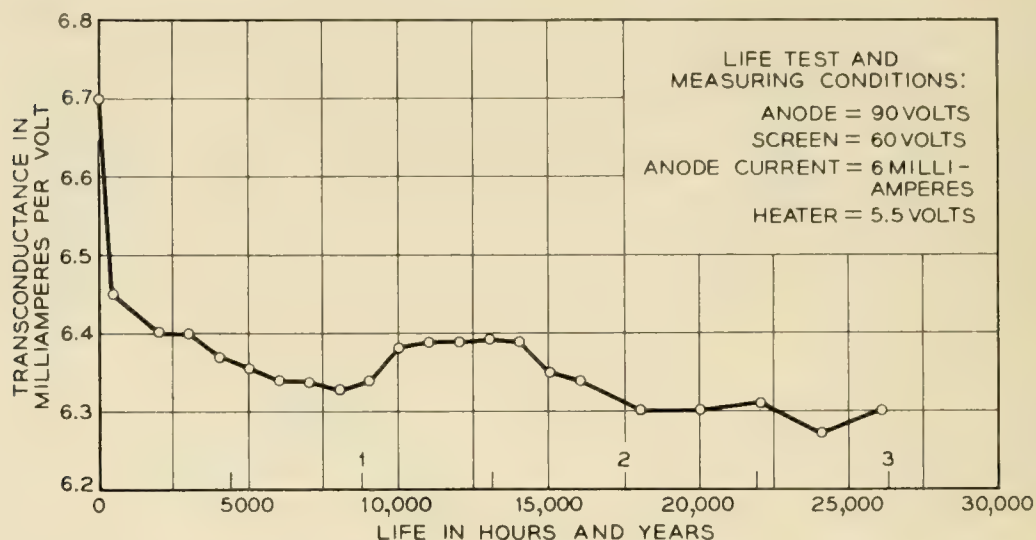


FIG. 13 — Behaviour of a group of 92 type 6P12 tubes over a period of three years.

In the early stages of the test there were eight mechanical failures. The cause in all instances was identified and corrected in subsequent production before the start of the T.A.T. project.

Mechanical Characteristics

The chief mechanical characteristics of the 6P12 have been mentioned before in that they are, as explained, very similar to those of the 6P10. A photograph of the interior of the tube is shown in Fig. 14.

Tube Selection Techniques

Not all the tubes, found after production to be potentially suitable for the British T.A.T. amplifier, remained equally suitable after the life test period of about 4,000 hours. A brief account of how the best were selected is given below.

The fact that every tube had to pass conventional static specification limits needs little emphasis here. This test was, however, supplemented by three additional types of specification. First, every tube was tested in a functional circuit, simulating that stage of the amplifier for which the tube was ultimately intended. Here measurements were made of shot noise (appropriate to first stage usage) and harmonic generation (appropriate to the output stage) in addition to the usual measurements of transconductance, anode impedance and working point.

Second, all tubes were subject to intensive visual scrutiny in which some 80 specific constructional details were checked for possible faulty assembly.

Third, the life characteristics of transconductance, total emission and working point were examined over the test period of about 4,000 hours for unsatisfactory trends. Although this type of specification is more difficult to define precisely, its application is probably more rigorous and exacting than any of the previous specifications.

Only if a tube passed the conventional test and the three supplementary tests was it considered adequate for inclusion in a repeater.

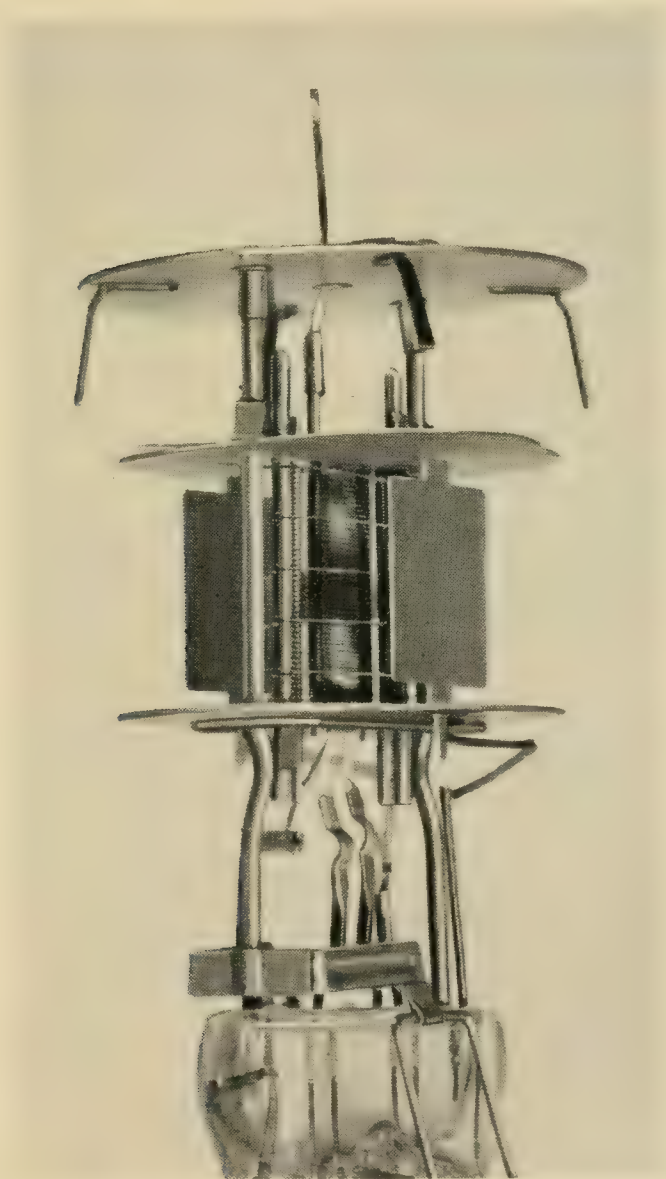


FIG. 14 — View of interior of a 6P12 type tube.

CONCLUSION

The laying of the present repeatered transatlantic cable represents by far the most ambitious use to date of long life, unattended electron tubes. On this project alone there are 390 tubes operating on the ocean bottom. If to this number are added the ocean bottom tubes from earlier shorter systems, those used in the Alaskan cable completed a few months ago, and those to be used in the California-Hawaiian cable to be laid in 1957, the total number on the ocean bottom will be about one thousand. The capital investment dependent on the satisfactory performance of these tubes is probably about one hundred million dollars — strong evidence of faith in the ability to produce reliable and trustworthy tubes.

It is of interest to note that the two groups working on the tubes on opposite sides of the Atlantic had no intimate knowledge of each other's work until after the tube designs had been well established. As a result of subsequent discussions, it has been surprising and gratifying to find how similarly the two groups look at the problems of reliability of tubes for submarine cables.

The authors would be completely remiss if they did not mention the contributions of others in the work just described. These projects would have been impossible if it were not for the enthusiastic, cooperative and careful efforts of many people working in varied fields. Over the years chemists, physicists, electrical and mechanical engineers, laboratory aides, shop supervisors and operators all have made essential contributions to the projects. It would be impractical and unfair to attempt to single out for mention the work of specific individuals whose contributions are outstanding. There are too many.

REFERENCES

1. K. G. Compton, A. Mendizza and S. M. Arnold, Filamentary Growths on Metal Surfaces — "Whiskers," *Corrosion*, **7**, pp. 327-334, Oct. 1951.
2. M. Benjamin, The Influence of Impurities in the Core-Metal on the Thermionic Emission from Oxide-Coated Nickel, *Phil. Mag. and Jl. of Science*, **20**, p. 1, July, 1935.
3. H. E. Kern and R. T. Lynch, Initial Emission and Life of a Planar-Type Diode as Related to the Effective Reducing Agent Content of the Cathode Nickel (abstract only), *Phys. Rev.*, **82**, p. 574, May 15, 1951.
4. G. H. Metson, S. Wagener, M. F. Holmes and M. R. Child, The Life of Oxide Cathodes in Modern Receiving Valves, *Proceedings I.E.E.*, **99**, Part III, p. 69, March, 1952.
5. G. H. Metson and M. F. Holmes, Deterioration of Valve Performance Due to Growth of Interface Resistance, *Post Office Electrical Engineers Journal*, **46**, p. 193, Jan., 1954.
6. M. R. Child, The Growth and Properties of Cathode Interface Layers in Receiving Valves, *Post Office Electrical Engineers Journal*, **44**, p. 176, Jan., 1952.
7. G. H. Metson, A Study of the Long Term Emission Behaviour of an Oxide Coated Valve, *Proceedings I.E.E.*, **102**, Part B, p. 657, Sept., 1955.

Cable Design and Manufacture for the Transatlantic Submarine Cable System

By A. W. LEBERT,* H. B. FISCHER* and M. C. BISKEBORN*

(Manuscript received September 19, 1956)

The transatlantic cable project required that two repeatered cables be laid in the deep-water crossing between Newfoundland and Scotland, and one across the shallower waters of Cabot Strait. The same structure was adopted for the cables laid in the two locations.

This paper discusses the considerations leading to design of the cable and describes the method of manufacture, the means and equipment for control of cable quality, the process and final inspection procedures, the electrical characteristics of the cable, and factors relating to mechanical and electrical reliability of the final product.

DESCRIPTION OF CABLE

General features of the cable structure adopted for the transatlantic cable project¹ are illustrated in Fig. 1. The cable consists of two basic parts: (1) the coaxial, or the electrical transmission path, and (2) the armor or outer protection and strength members.

The coaxial is made up of three parts: (1) the central conductor, (2) the insulation, and (3) the outer or return conductor. The central conductor is composed of a copper center wire surrounded by three helically applied copper tapes. The insulation is a polyethylene compound which is extruded tightly over the central conductor. The insulated central conductor is called the cable core. The outer or return conductor is composed of six copper tapes applied helically over the insulation.

The protection and strength components shown in Fig. 1 for the type D deep water cable are provided by a teredo tape of thin copper applied over the outer coaxial conductor, a fabric tape binding, a layer of jute rove for armor bedding, the textile covered armor wires and finally, two layers of jute yarn flooded with an asphaltum-tar compound. This cable is characterized by the extra tensile strength of its armor wires and by the extra precautions taken to minimize corrosion of these wires.

* Bell Telephone Laboratories.

At the shallow water shore ends, the armor types are characterized by the use of mild steel wires which are increased in diameter in steps as the landing is approached. These types are designated A and B and will be described in greater detail later.

The transmission loss of this cable structure at the top operating frequency of 164 kc is 1.6 db per nautical mile and is 0.6 db per nautical mile at 20 kc, which is the lower end of the frequency band. The high frequency impedance of the cable is about 54 ohms.

BASIS OF DESIGN

A coaxial structure was first used for telephone and telegraph service in a submarine installation in 1921, between Key West and Havana, Cuba. Three coaxial cables with continuous magnetic loading and no submerged repeaters were laid. One telephone circuit and two telegraph circuits were provided in each cable for each direction of transmission.

In 1950, a pair of submarine coaxial cables,² which included flexible submerged repeaters, was laid between Key West and Havana, Cuba. Each cable furnished 24 voice circuits. One cable served as the "go" and the other as the "return" for the telephone conversations. The transatlantic telephone cable design is similar to this cable except that the nominal diameter of the insulation is 0.620" instead of 0.460". An outstanding difference between the transatlantic and Key West-Havana systems is cable length — about 2,000 nautical miles as compared with 125. This difference influenced significantly the permissible electrical and mechanical tolerances applying to the cable structure.

The installation of some 1,200 miles of cable with island based repeaters for a communication and data transmission system for the U. S. Air Force, between Florida and Puerto Rico,³ followed the 1950 submarine cable system. The design of this cable is identical with that of the transatlantic cable, except for differences in the permissible dimensional tolerances on the components of the electrical transmission path. Data obtained on the electrical performance of the Air Force cable provided the transmission characteristic to which the repeaters for the transatlantic project were designed.

The design of this cable installation was the result of many years of cable development effort, which was guided by the successes and failures of the earlier submarine telegraph cables. The 1950 Key West-Havana and the Air Force cables differed from the earlier structures in one important respect, namely, the lay of the major components. A series of fundamental design studies during the 1930's and 1940's and extensive field tests in the Bahamas in 1948 demonstrated that having

the same direction of lay of the major components of a cable was very important in minimizing kinking and knuckling. In addition, other laboratory tests pointed the direction for the adoption of new materials and techniques in the manufacture of these cables. These and subsequent improvements in materials and manufacturing techniques were included in the transatlantic cable design.

Since the electrical characteristics of a cable have a direct bearing on the overall system design and performance, considerable emphasis was placed on this phase of the design of the transatlantic cable. The size of

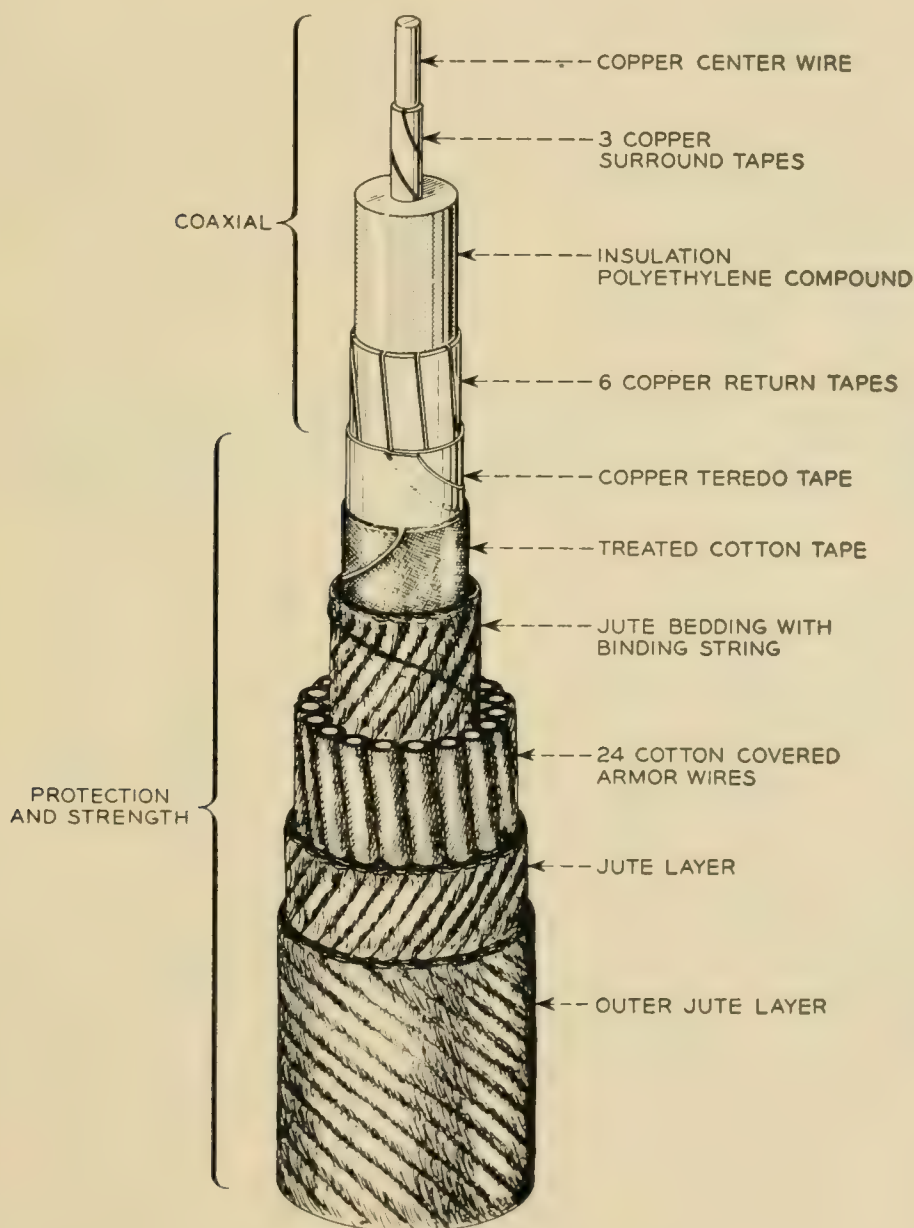
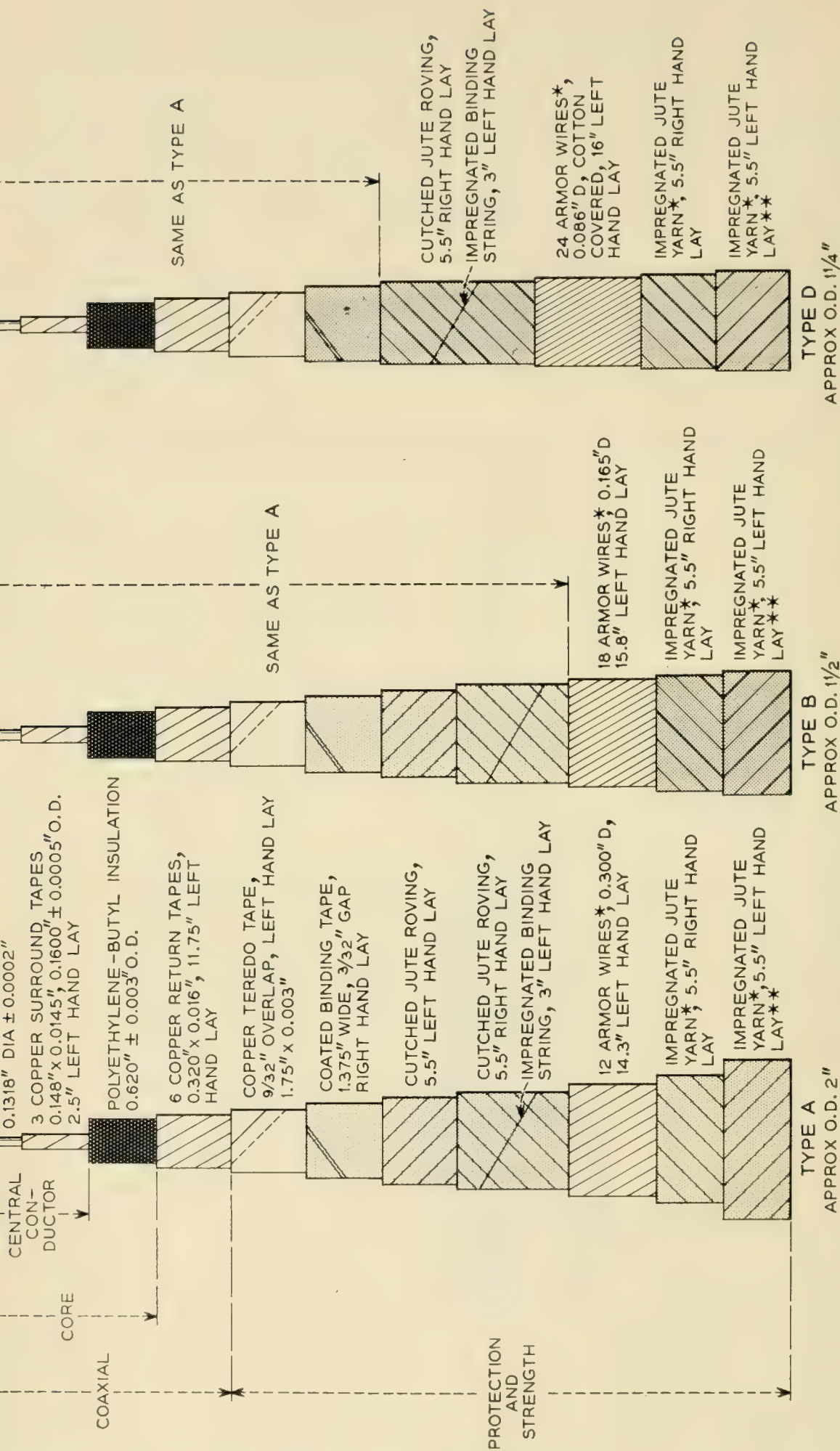


Fig. 1 — Structural features of the deep water type of cable.



*HOT FLOODING COMPOUND APPLIED DURING THESE OPERATIONS
 **COATED WITH CALCIUM CARBONATE

Fig. 2 — Design details of the types of transatlantic telephone cable for different depths of water.

central conductor and core used in the Air Force cable resulted in low unit attenuation and low dc resistance. These advantages resulted in the adoption of this size of cable. However, the outside diameter of the core and the diameter of the central conductor of the coaxial do not fulfill the requirements generally described as optimum for minimum attenuation. Mathematical analysis shows that there is a preferred diameter ratio which results in minimum transmission loss. For the 0.620" core diameter employed in the Air Force cable, the central conductor diameter chosen was smaller than the ideal central conductor required to satisfy the preferred diameter ratio. The diameter chosen retained the central conductor size which the Key West-Havana and Air Force cables proved to be satisfactory from a manufacturing standpoint. The choice was also compatible with the dc resistance requirement for transmission of power over the cable to each of the repeaters.

While production of the cable was proceeding, cable manufactured to the transatlantic specification was tested near Gibraltar in March, 1955. These tests provided a final evaluation of the mechanical and electrical characteristics of this cable before the actual laying of the transatlantic link.

DETAILS OF STRUCTURAL DESIGN OF CABLE

The structural features of the coaxial and of types A, B and D armor are summarized in Fig. 2.

A composite central conductor was chosen to provide a conductive bridge across a possible break in any one of its elements, due to a hidden defect, such as an inclusion of foreign material in the copper. The dimensions of the components of the central conductor were precisely controlled, and a light draw through a precision die was used to compact and size the assembly.

Use of high molecular weight polyethylene (grade 0.3) for core insulation is a major departure from the materials used in early submarine telephone and telegraph cables. The development of synthetic polymers such as polyethylene had led to the replacement of gutta percha as cable insulation, since polyethylene possesses better dielectric properties and mechanical characteristics and is lighter in weight.

Ordinary low molecular weight polyethylene is subject to environmental cracking, especially in the presence of soaps, detergents and certain oils. High molecular weight material is much less subject to cracking, and by adding 5 per cent butyl rubber, further improvement in crack resistance is obtained.

Six copper tapes applied helically over the core comprised the return

conductor and thus completed the coaxial structure. The dimensions of these tapes were precisely controlled. The helical structure was chosen to impart flexibility to the coaxial.

Insulation of some of the early submarine telegraph cables suffered from attack by marine borers such as the teredo, pholads and limnoria. To protect against such attack, a thin metallic tape was placed over the insulation in the early submarine cables. The necessity for such protection for the transatlantic cables, especially in deep water, may be questioned, but the moderate cost of this protection was considered cheap insurance against trouble. The copper teredo tape was applied directly over the return conductor, as a helical serving with overlapped edges to completely seal the coaxial from attack by all but the smallest marine organisms.

A cotton tape treated with rubber and asphaltum-tar compound was applied over the teredo tape to impart mechanical stability to the coaxial during manufacture. A small gap between adjacent turns of the helix was specified to permit ready access of water to the return tape structure and to the surface of the core. The use of a gap was based on laboratory tests which showed that transmission loss was dependent to a modest extent on thorough wetting of the exterior of the coaxial. Since transmission loss measurements are made on repeater sections of cable shortly after manufacture to determine whether any length adjustments are required, it was essential that the wetting action be as rapid as possible.

The design of the protection and strength components of the cable was modified according to the depth of the water in which the cable was to be laid. To prevent damage to the coaxial by any cutting action of the armor wires during manufacture and laying, a resilient cushion of jute roving was placed between the armor wires and coaxial. For type-D cable, a single layer of jute was used; for types A and B cable, the bedding was made up of two layers of jute. To protect this jute from microbiological attack, a cutting treatment was employed. The traditional

TABLE I

Type	Armor Wire			Application
	Number of Wires	Diameter in Inches	Material	
A	12	0.300	Mild Steel	Up to 350 fath.
B	18	0.165	Mild Steel	350 to 700 fath.
D	24	0.086	High Strength Steel	Greater than 700 fath.

cutching process consists of treating the jute with a vegetable compound called catechu or cutch.

A armor wires were applied over the bedding jute. The use of heavy or intermediate weight near shore has been established by experience with ocean cable. This type of armor is generally employed where the cable may be exposed to wave action, bottom currents, rocks, icebergs, ship's anchors and fishing trawlers. A lighter weight structure having higher tensile armor wires is needed in deep water. Table I shows the essential differences between the armor types employed in the transatlantic cable and the approximate range of depths in application.

In addition to the above armor types, a shore length of 0.6 nautical mile was provided with an insulated lead sheath under Type A armor to facilitate preferred grounding arrangements and to provide signal to noise improvement.

Where the tensile strength of the armor wires is most important, as in the type D design, each of the wires was protected against corrosion by a zinc galvanize plus a knitted cotton serving or helically applied tape, the whole assembly being thoroughly saturated with an asphaltum-tar compound. The effectiveness of such protection is clearly apparent when early submarine cables, which used this protection, are recovered and examined. For the heavier armor types, the protection was similar to that of type D, except that the textile serving was replaced by a dip treatment to coat each wire with an asphaltic compound.

As the armor wires were applied to each type of cable additional protection was obtained by flooding the cable with a special asphaltum-tar compound and then applying two layers of jute yarn over the wires. The jute yarn was impregnated with an asphaltum-tar compound before application to the cable and then flooded with another asphaltum-tar compound after application. Formulation of cable flooding materials required the use of compounds having a relatively high coefficient of friction to avoid slippage of the cable on the ship's drum during laying.

To assure satisfactory handling characteristics during the laying operation, all of the metallic elements of the cable were applied with a left-hand direction of lay and the lengths of lay (except for the teredo tape) were chosen so that approximately the same helical length of material was used per unit length of cable. Since the teredo tape was relatively soft and ductile compared to that of the other metallic components, it was not necessary to equate its helical length to that of the other components. Width and lay of the teredo tape were selected to give a smooth, tight covering.

The choice of direction and length of lay of the jute layers was based

on experience with cable in factory handling and laying trials. Experience, particularly with the direction of lay, has shown that improper choice of lays for the two outer layers of jute may result in a cable that is difficult to coil satisfactorily in factory and ship storage tanks. The combination of lays selected for the cable components provided good performance in all the handling operations, including the final laying across the Atlantic.

MANUFACTURE OF THE CABLE

Before considering the manufacture of the cable, it should be understood that the repeater gain characteristic was designed to compensate for the loss characteristic of the cable. Therefore, once this loss characteristic was established, it was essential that all cable manufactured conform with this characteristic.

To obtain the required high degree of conformance, close control had to be kept over all stages of manufacture and over the raw materials. Controls to guide the manufacture of the cable were set up with two broad objectives:

1. To produce a structure capable of meeting stringent transmission requirements.
2. To assure that the manufactured cable could be laid successfully and would not be materially affected by the ocean bottom environment for the expected life of the cable system.

Attainment of a final product capable of meeting the stringent transmission requirements is described in a later section of this paper. Process and raw material controls in manufacture were provided by an inspection team which checked the quality of the various raw materials and the functioning of the several processes during the manufacturing operations. This type of inspection coverage is somewhat unique with submarine cable. It assures the desired final quality by permitting each error or accident to be investigated and corrected on an individual basis.

Cable for the transatlantic crossing was manufactured in America by the Simplex Wire and Cable Company and in England by Submarine Cables, Limited. Differences in machinery and equipment in the plants of the two manufacturers necessitated minor differences in the sequence of the operations and in the processes. The sequence of operations in assembly of the cable was as follows:

Step No.	Operation
1	Stranding of central conductor
2	Extrusion of insulation
3	Runover examination, repair where necessary

- 4 Panning and testing of core
- 5 Jointing of core
- 6 Application of return tapes, teredo tape, fabric tape, jute bedding and binding string
- 7 Application of armor wire and outer jute layers
- 8 Storage in tanks, testing
- 9 Splicing in repeaters, testing

The only important difference in the sequence of the manufacturing operations at the two plants was the use of separate operations for steps 6 and 7 at Submarine Cables and the combination of these operations in one machine at Simplex.

Other minor differences in process methods related to raw materials. For example, the American supplier purchased polyethylene already compounded with butyl rubber and antioxidant in granule form, ready for use. The British supplier purchased polyethylene, butyl rubber and antioxidant separately, and performed the compounding in the cable factory.

STRANDING OF CENTRAL CONDUCTOR

The central conductor was stranded on a machine which included a revolving carriage with suitable arbors for the three surround tapes. It was equipped with brakes designed to assure equal pay-off tension among the tapes and with detectors to automatically stop the strander in case of a tape break. Each tape was guided through contoured forming rolls to shape the tape to the center wire.

The joints between successive reels of wire and tape used in fabrication of the central conductor were butt brazed. The brazes were staggered to avoid more than one braze in a given cross-section of the conductor. The quality of the brazes in these components was controlled by a qualification technique described below in the section on core jointing.

The strand was drawn through several forming dies to size the finished diameter of the central conductor accurately. No lubrication was used because the removal of the resultant residues, which could contaminate the polyethylene insulation, was difficult. The taper in the central conductor diameter due to the die wear was controlled by appropriate replacement of tungsten carbide dies, where used, or by the use of a diamond die where the rate of die wear is less than 1 or 2 micro-inches per mile.

The stranding area in both plants was enclosed and pressurized to guard against dirt and dust settling on the central conductor. A high standard of cleanliness was maintained for parts of the machine which touched the conductor or its components. Undue wear of the guide faces

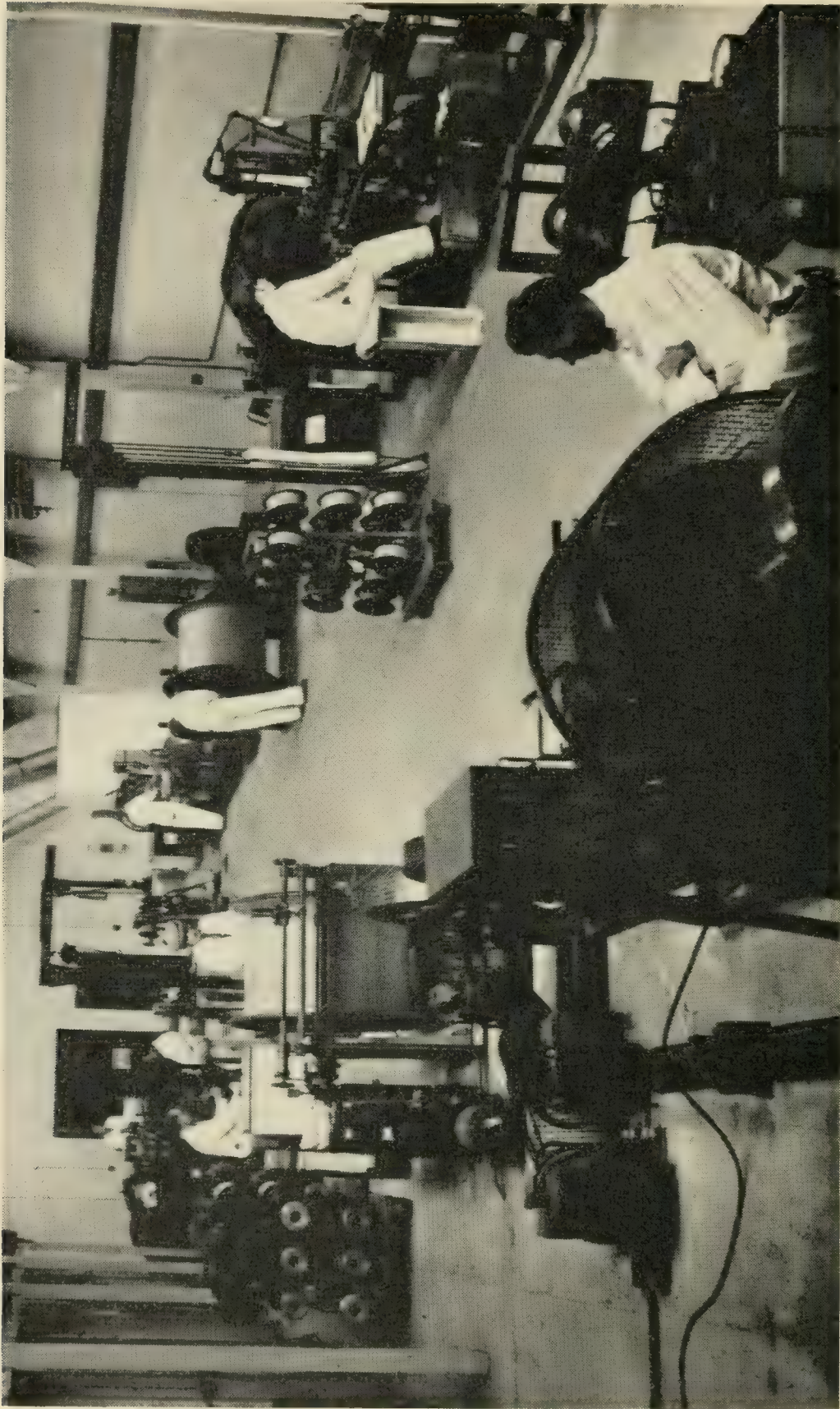


Fig. 3. — View of typical manufacturing area for stranding of central conductor.

or capstan sheaves was cause for replacement of the sheaves and adjustment of the machine. A photograph of the stranding area is shown in Fig. 3.

EXTRUSION OF INSULATION

To avoid possible contamination of the polyethylene insulating compound in the extruder-hopper loading area, a pressurized enclosure prevented entry of air-borne dust and dirt, and the containers of polyethylene compound were cleaned with a vacuum cleaner before being brought into the hopper area. A fine screen pack placed in the extruder reduced the possibility of contamination in the core.

In passing through the extruder, the central conductor was payed-off of a large reel with controlled tension, into the pay-out capstan, through an induction heater, through a vacuum chamber, and thence into the cross-head of the extruder. The induction unit heated the central conductor and provided means for controlling the shrinkback of the core insulation and the adhesion of the conductor to the insulation. Shrinkback is a measure of the contained stresses in the insulation.

On the output side of the extruder, the core was cooled in a long sectionalized trough containing progressively cooler water from the input to the output end. The annealing of the polyethylene in the cooling trough also served to hold the shrinkback of the core to a low value. The extrusion shop is shown in Fig. 4.

An important addition to the extrusion operation consisted of the use of an improved servo system to control the extruder automatically to attain constant capacitance per unit-length of coaxial. The system used is described in a subsequent section.

RUNOVER EXAMINATION

Following extrusion, the core was subjected to continuous visual and tactual examination in a rereeling operation called "runover". The purpose of the runover operation was the detection of inclusions of foreign material in the dielectric and the presence of abnormally large or small core diameters. Core not meeting specification requirements was cut out or repaired.

In addition to visual inspection of the cable, examination of short lengths of core was made at regular intervals with a shielded source of light arranged to illuminate the interior of the dielectric material. Provision of this internal illumination facilitated detection of particles of foreign material well beneath the surface of the core. Strip chart records

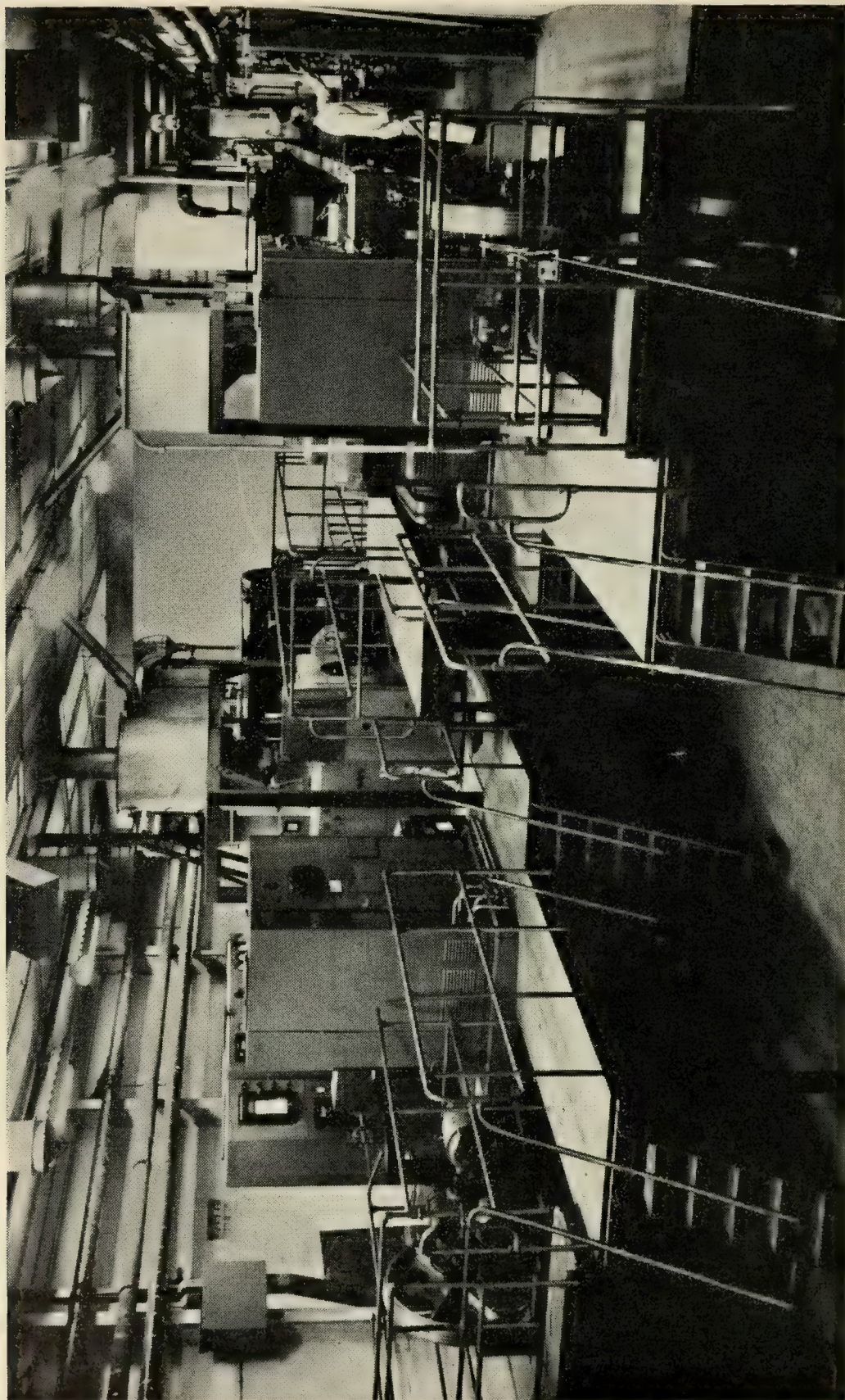


Fig. 4 — View of typical manufacturing area for extrusion of core insulation.

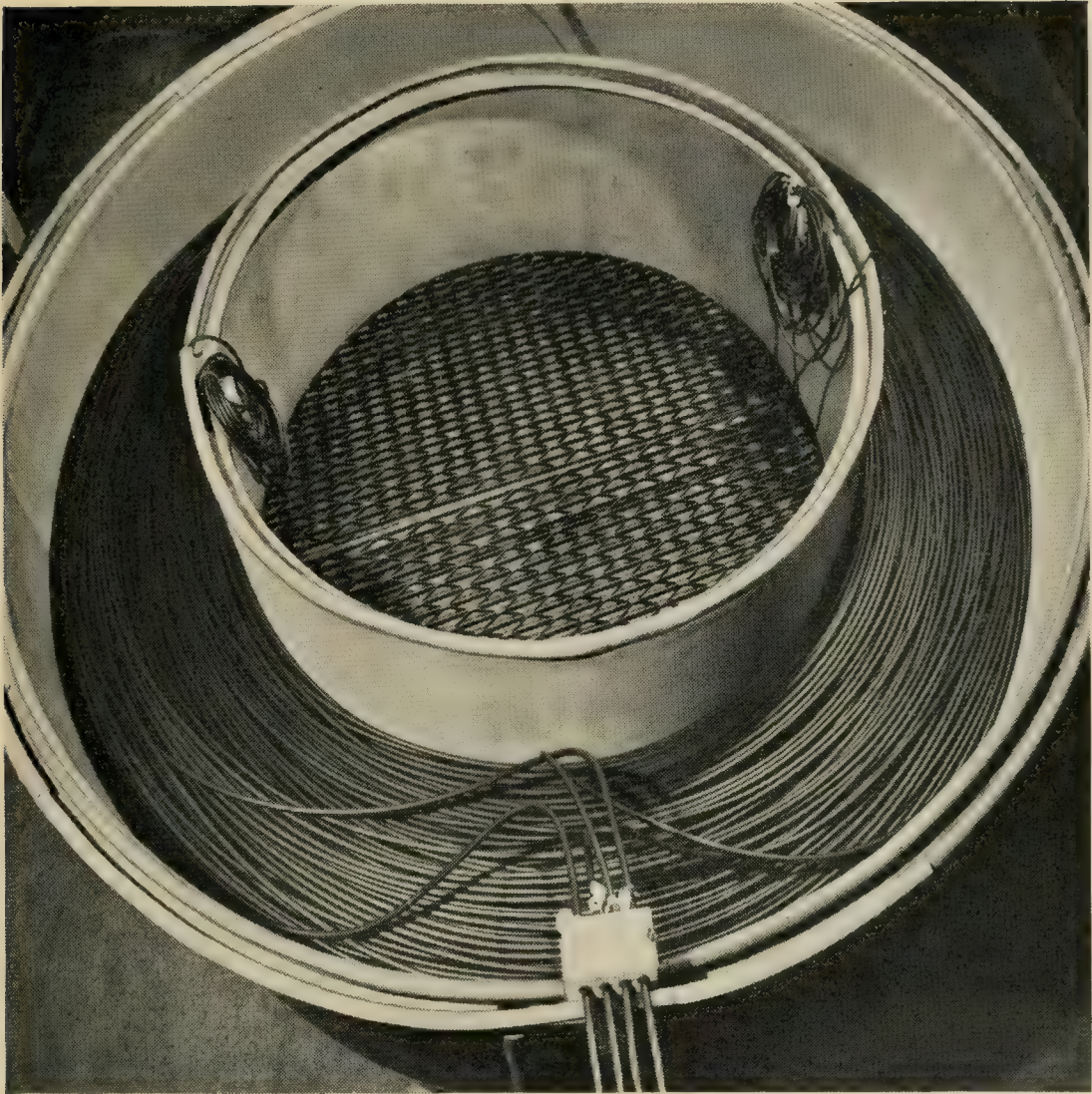


Fig. 5 — Special water tank or pan for tests on immersed cable core.

of the unit-length capacitance of the core obtained during extrusion were used as a guide in searching out regions of uncertainty.

PANNING AND TESTING

After runover, the core was coiled in tanks of water, as shown in Fig. 5. Precautions were taken to remove the air dissolved in the water and thus prevent the formation of bubbles on the surface of the core. The water was also temperature controlled and circulated to maintain uniform temperature throughout the tank. Thermocouples placed at different levels in the tank determined when the temperature was uniform. Measurements of dc conductor resistance, ac capacitance, insulation re-

sistance, and dielectric strength were then made. These measurements are discussed in detail in a later section of this paper.

JOINTING OF CORE LENGTHS

The cable core was manufactured in lengths much shorter than a repeater section, which necessitated connecting the individual core lengths together. Jointing techniques consisted of brazing the central conductor with a vee-notch type of junction and of molding in a short section of the polyethylene insulation. After silver-soldering the vee-joint, a safety wire consisting of four fine gauge tinned copper wires was bridged across the junction in an open helix and soft soldered at the ends. Extreme care was taken to remove any excess rosin and to eliminate any sharp points on the ends of the safety wires. The safety wire is intended to maintain continuity of the electrical path in case the braze should fail.

Visual examination of brazes in the actual cable was the only means available for their final inspection. To assure a high degree of quality on these brazes, a system was devised for checking the performance of the operator and the brazing machine initially and at frequent intervals through the use of sample brazes in each of the components, which were tested to destruction. To control the uniformity of brazes, the brazing of the copper wire and tapes was made as automatic as possible by the use of controlled pressure on the components, appropriate sized wafers of silver solder, and an automatically timed heat cycle. The tests on the brazes used to qualify operators and machinery indicated that a high degree of braze performance was achieved.

Pressure and temperature were carefully controlled during the molding of the insulation over the conductor. Periodic checks similar to those described for brazes were made on operator and molding-machine to maintain a satisfactory level of performance. In addition, each molded joint placed in the actual cable was X-rayed and subjected to a high voltage test while immersed in water.

APPLICATION OF RETURN TAPES AND ARMOR WIRES

After the core lengths were joined together, they were pulled through the return taping and armoring operations. The machine for applying return tapes was designed specifically for the purpose and was similar in characteristics to the corresponding portion of the strander for the central conductor. Controlled pay-off tension, automatic breakage detectors, precision guides, and contoured forming rolls to shape the tape, were incorporated in the construction.

The return tape, teredo tape, and fabric tape were applied from taping heads, and the bedding jute and binding string were applied from serving heads in a tandem operation. Another set of tandem operations included the application of armor wires, outer jute layers and the appropriate asphaltum-tar flooding compounds. In the American suppliers plant, both sets of tandem operations were combined into one continuous production line. In the British plant, these operations were divided into two separate production lines. A view of the armoring machine area is shown in Fig. 6. Following the application of the flooding compounds, whiting is applied either at the take-up capstan on the armoring line or in the storage tanks as the cable is coiled.

To avoid core damage, the flow of hot flooding compound was stopped when the cable in the armoring line was stopped. One of the major sources of such stoppages was the reloading of the various applying heads.

STORAGE AND TESTING

In a continuous haul-off operation, the cable was conveyed from the armoring machine to the tank house for storage. The cable was coiled in spiral layers, called flakes. Each flake started at the outside rim of the tank and worked toward the central cone. Several 37-nautical mile repeater sections were stored in each tank.

Water was circulated through the cable tanks to establish uniform temperature conditions throughout the mass of cable. When thermocouples located at appropriate points in the tank indicated that the cable temperature was uniform, measurements were made of attenuation, internal impedance irregularities and terminal impedances, dc resistance, dc capacitance, insulation resistance, and dielectric strength.

To facilitate these tests without interrupting production, successive repeater section lengths were placed in alternate tanks. By this procedure, a group of four or five sequential repeater section lengths, called an "ocean block", was stored in two tanks. The ends of each repeater section were brought out of the tank to a splicing location. After all tests were completed and the specification requirements met, the repeaters were spliced in. Testing of the ocean block for transmission performance completed the manufacturing operations.

RAW MATERIALS

Stringent requirements were placed on all raw materials used in the manufacture of the transatlantic cable. Detailed specifications covered

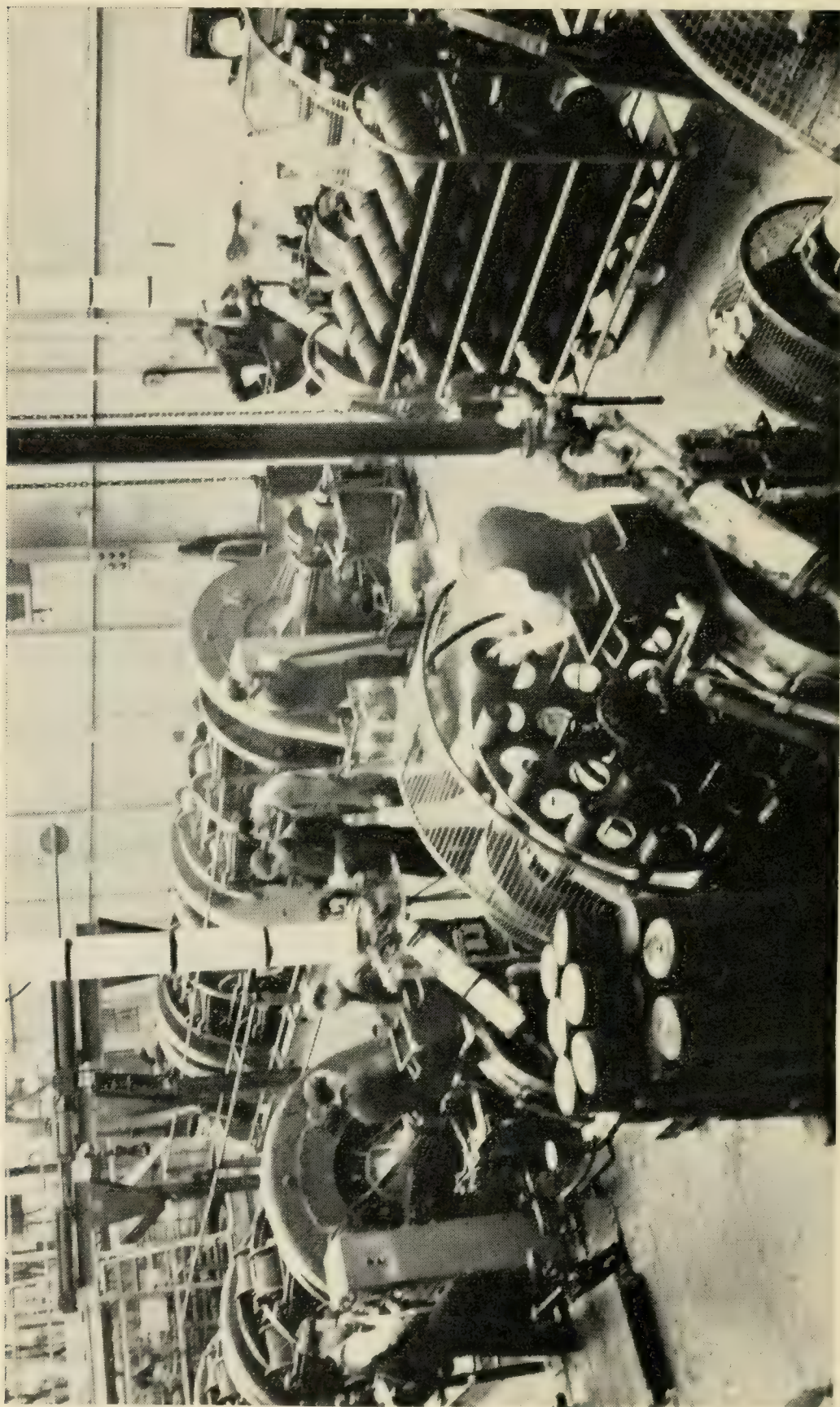


Fig. 6 — View of shop for application of cable armor.

the basic requirements and the methods of controlling their quality on a sampling inspection procedure. The requirements for the materials were established to insure that their use would not jeopardize the life of the cable. Since cable life is critically related to the integrity of the insulation, all materials had to be scrutinized for their tendency to cause environmental cracking. These tests were necessarily made on an accelerated basis. Since no correlation exists at present between accelerated tests and long term (20 year) life tests, only conservative design selections can be justified.

Close tolerances such as ± 0.0002 inch for the diameter of the solid center wire in the central conductor were specified for all copper components of the coaxial. In addition, these components had to be free from slag or other inclusions, and the wire drawing and rolling of the tape had to be controlled to assure smooth surfaces, edges of prescribed shape, and freedom from filamentary imperfections. Compounds used in drawing and rolling operations were selected to minimize the possibility of contaminating or causing cracking of the polyethylene. Residual quantities of compound on the wire or tapes were removed prior to annealing, which was controlled to prevent the formation of oxides and to assure clean and bright copper.

The dielectric constant range of the polyethylene-butyl compound was limited to 2.25 to 2.29. These limits were determined by the limited accuracy of the measuring equipment available at the time. Restrictions covered the allowable amount of contamination since its presence in other than minute quantities might reduce the dielectric strength or degrade the power factor of the compound.

In addition, the melt index (a factor related to molecular weight) of the final insulating compound composed of polyethylene resin, butyl rubber, and antioxidant, was held to 0.15 to 0.50. The melt index of ordinary polyethylene used for insulation generally, is 2.0 or higher. Choice of the low index assured the maximum resistance to environmental cracking.

The cutting and fixing processes used in the manufacture of bedding jute were adjusted to limit the alkalinity of the jute because of the adverse effect of alkaline materials on polyethylene compound. Oils used in the spinning of the jute were selected to obtain types which were not strong cracking agents for polyethylene, and the quantities used were reduced to the workable minimum. The presence of such impurities as bark and roots was restricted to provide the desired fiber strength. The impregnation of the jute was controlled to ensure adequate distribution

of the coal tar throughout the fibers without having an excess that would make the jute difficult to handle during the armoring process.

The size, composition, and processing of the armor wires were also placed under close control. Purity, tensile strength, and twist requirements were designed to ensure that the wire could be applied to the cable, and welded, and that it could withstand the expected tensions during laying and pickup. Strength considerations made it mandatory that inclusions of slag or piping of the wire be eliminated. Piping is an unusual condition encountered during rolling or drawing which results in a hollow shell of steel which may be filled with slag.

CONTROL OF TRANSMISSION CHARACTERISTICS

From a broad point of view, the attainment of a final product capable of meeting the stringent transmission requirements was achieved by the following basic steps.

1. Precision control of the dimensions of the copper conductors, including the diameter of the fabricated central conductor.

2. Automatic control of the insulating process to maintain a constant capacitance, thus compensating for deviations in central conductor diameter and dielectric constant of the insulation.

3. Factory process control, by means of a running average of the measured attenuation characteristic of current production, to guide the adjustment of suitable parameters when necessary.

As indicated in the sections above on the method of manufacture and the control of raw materials, precautions were taken to obtain a central conductor that had predictable electrical performance, and a controlled taper in overall diameter along its length. The need for such effort is explained by consideration of the factors that determine the attenuation of a coaxial structure.

The attenuation, α , of the cable is directly proportional to the ac resistance, R , and inversely proportional to the characteristic impedance, Z_0 , as a satisfactory approximation. That is,

$$\alpha = \frac{aR}{Z_0} db/nm$$

where "a" is a coefficient depending on the units. It is thus clear that control of α may be attained by control of R and Z_0 . Since the resistance is a function largely of the diameter of the central conductor, and since it is held to close tolerances, the constancy of impedance completes the requirement for attenuation control.

The characteristic impedance of a transmission line is determined by:

$$Z_0 = b \sqrt{\epsilon} \log \frac{D}{d} \text{ ohms}$$

where ϵ is the dielectric constant of the insulating material, D is the inside diameter of the outer conductor, d is the diameter of the inner conductor, and b is a numerical coefficient. If the dielectric constant of the insulating material (polyethylene) does not vary, control of characteristic impedance reduces to control of capacitance. This follows from the fact that the capacitance, C , is related to the D/d ratio as follows:

$$\frac{D}{d} = \text{antilog} \frac{k\epsilon}{C}$$

where k is a numerical coefficient.

Precision control of capacitance during the insulating process is achieved by a double-loop linear servo system, as shown in a simplified block diagram, Fig. 7. The two loops consist, respectively, of one capable of introducing relatively fast capacitance corrections of only modest accuracy and of one capable of highly precise capacitance control on a relatively long time basis. The servo system controls the capstan payout

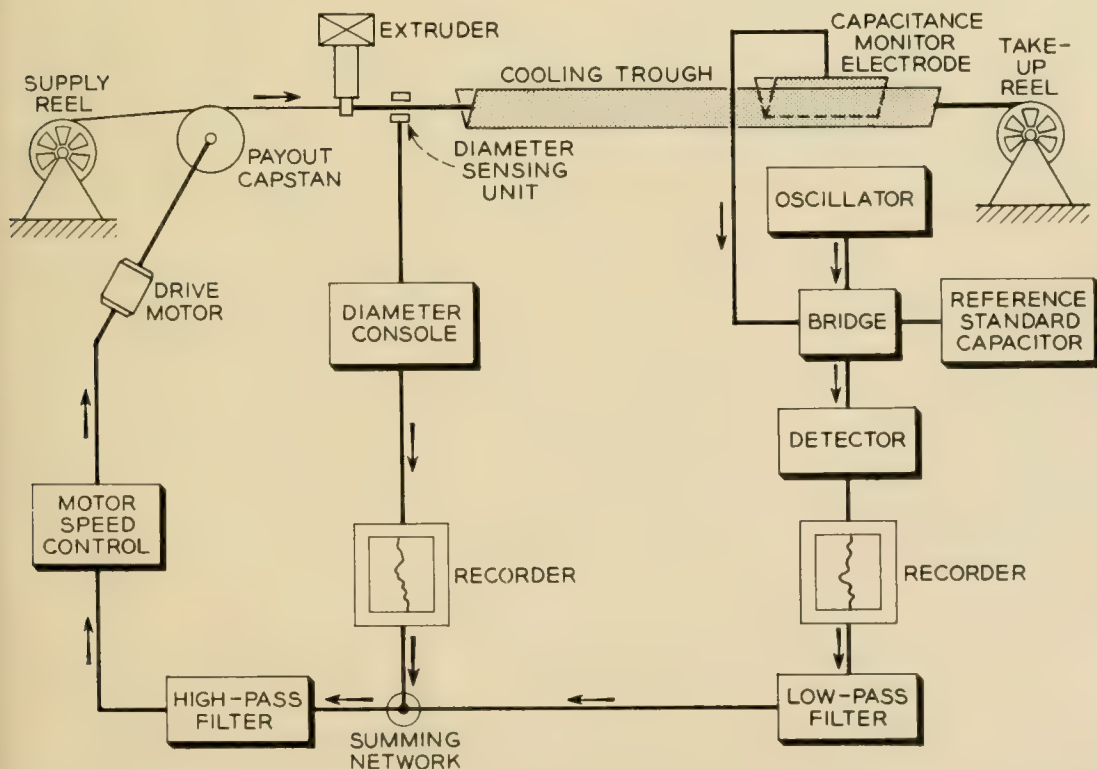


Fig. 7 — Simplified block diagram of capacitance monitor servocontrol system.

speed of the central conductor feeding into the extruder applying the core insulation, as shown in the block diagram. Since the extruder delivers insulating material at a constant rate, an increase in central conductor speed results in a thinner than normal wall of insulation and thus causes an increase in capacitance.

The sensing element for the control loop consists of a capacitance monitor. This is a device capable of measuring the unit length coaxial capacitance of the cable core continuously as it moves through the water in the trough. Since the capacitance of a polyethylene insulated core is temperature sensitive, the monitoring electrode must be located at a point in the cooling trough where the temperature of the core is stable and known to a degree commensurate with the overall accuracy objectives. The distance from the extruder to the electrode corresponds to about 10 minutes of cooling time; hence, a servo system based on this loop would be necessarily slow, due to the 10-minute delay in detecting a drift in capacitance.

Analysis shows that fast capacitance information of only moderate accuracy may be used in combination with the slow loop to speed up the response of the overall system to a satisfactory degree, without sacrifice of precision of the slow loop. The sensing element used for the fast loop consisted of a light-ray diameter gauge, which measures the diameter (changes in diameter are the approximate inverse of the capacitance) of the hot core close to the extruder. The slow and fast data are combined to control the extruder, as shown in the block diagram.

The servo constants were chosen to minimize the deviations in unit length capacitance occurring in core lengths corresponding to less than $\frac{1}{4}$ wave length of the top operating frequency. Stated in other words, the objective for choice of servo loop constants was to assure equality in the capacitance of all $\frac{1}{4}$ wave sections of core. Echo measurements indicated that a highly satisfactory degree of control was achieved. Overall servo system performance was such that the standard deviation of the capacitance of the core lengths manufactured for the two crossings was ± 0.1 per cent. The capacitance monitor electrode and the servo console is illustrated in Fig. 8.

ADJUSTMENT OF CONCENTRICITY

Means for setup and adjustment of the extrusion process to achieve relatively accurate centering of the conductor in its sheath of insulation was provided by a device called a concentricity gauge. This device operates on the principle that two small, plane electrodes on opposite

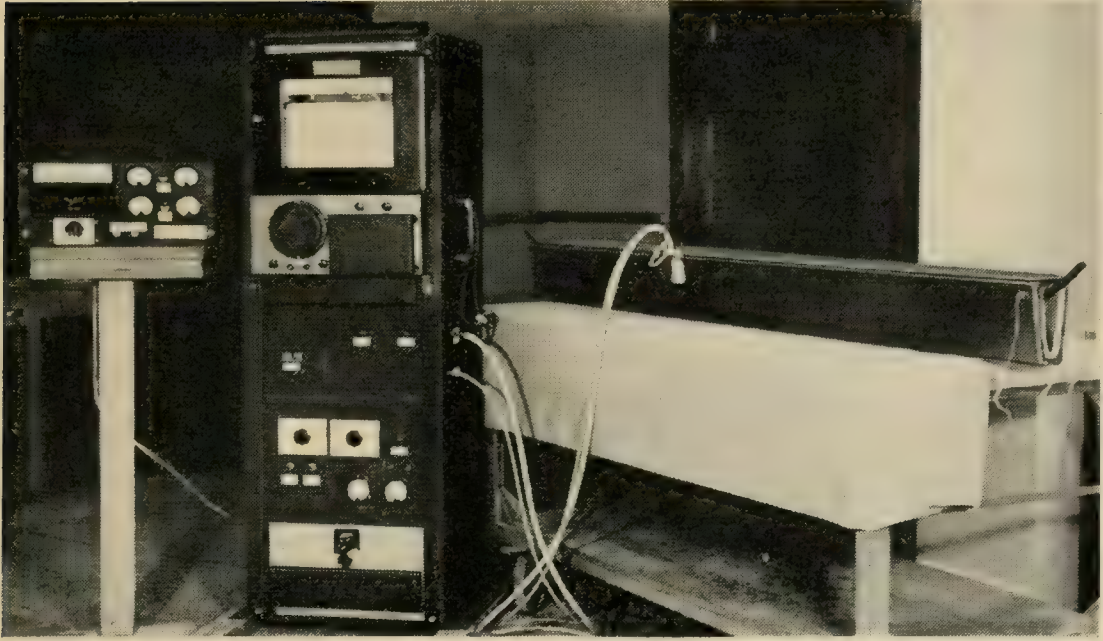


Fig. 8 — Photograph of capacitance monitor electrode and servo-controller console in laboratory setup.

sides of the core will have different direct capacitances to the central conductor, when the conductor is not properly centered.

A simplified block diagram of the concentricity gauge is shown in Fig. 9. Data obtained with two sets of electrodes displaced 90° were recorded on a strip chart recorder, with the output of the two sets of electrodes being displayed alternately. A satisfactory degree of centering was moderately easy to maintain.

ELECTRICAL MEASUREMENTS

To assist in achieving the goal of matching the cable and the repeater characteristics with a minimum of deviations, electrical measurements were made throughout the process and close tolerances were placed on the electrical parameters in each stage of production. Measurements on the repeater section lengths of cable were used as a final check to determine the extent to which all of the controls had been successful. The primary standards used were calibrated by the Bureau of Standards in the United States or the National Physical Laboratories in England. These precision standards were used to calibrate the bridges frequently.

The dc resistance of the central conductor was measured under constant temperature conditions with a precision type of Wheatstone bridge. The permissible range of resistance was 2.514 and 2.573 ohms

per nautical mile at 75°F. In practice, the spread of resistance values was well within these limits.

The 20-cycle core capacitance was also measured under constant-temperature conditions. For this measurement, the two ends of the central conductor were connected together and the measurements made between the central conductor and ground, which was provided by the water. A capacitance-conductance bridge was used for this purpose. The capacitance limits set initially were from 0.1726 to 0.1740 microfarads per mile at 75°F. Analysis of the core measurements indicates that at each factory the range of capacitance was held more closely than indicated, which illustrates the benefits of servo control to the insulating process.

The dc insulation resistance of the core was also measured by applying 500 volts for one minute. A minimum insulation resistance requirement of 100,000 megohm-miles at 75°F was established, but any lengths that had less than 500,000 megohm-miles were scrutinized for possible sources of trouble and were subject to rejection. As a general rule, insulation resistances considerably in excess of 500,000 megohm-miles were obtained.

The core was tested also at a voltage of 90,000 volts dc for a period of one minute. This test was designed to catch any gross faults in the

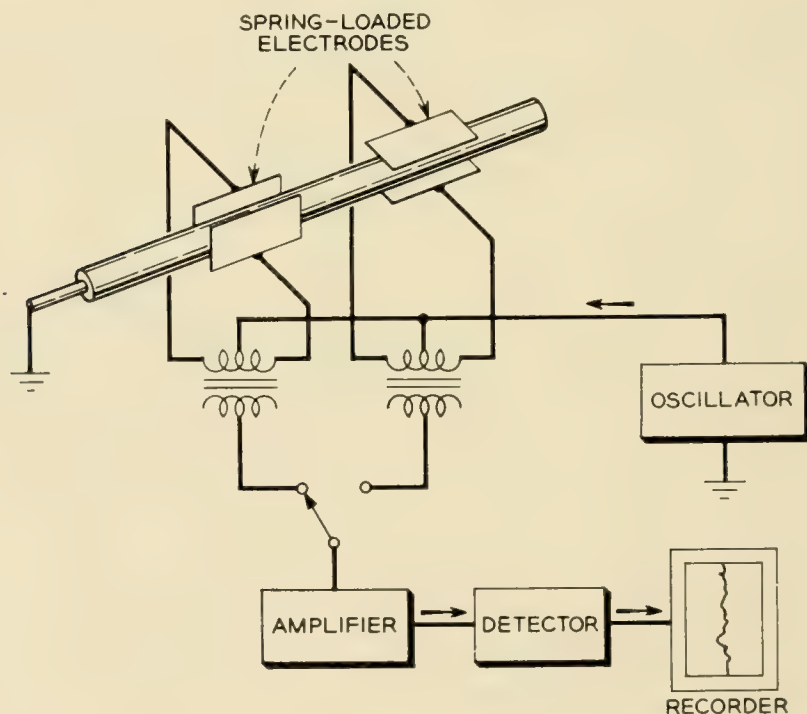


Fig. 9 — Simplified block diagram of concentricity gauge for continuous measurement of centering of central conductor.

core caused by foreign particles which escaped detection by the other mechanical and electrical tests made on the core.

As discussed under the section on jointing, the core lengths were assembled and joined together to form a repeater section of cable. In general, an effort was made to produce the core for a repeater section of cable on a particular strander, extruder, and armoring line, and to join the lengths together in the order of manufacture. Practical difficulties such as the fact that the outputs of two stranders were required to supply one extruder made it impossible to achieve this objective in all cases.

Capacitance deviations from the desired nominal resulted from a variety of causes, such as inaccurate control of the temperature of the water in the core cooling troughs and improper adjustment of the control apparatus. To minimize the reflection which would result from joining together two lengths of core of widely different capacitances, cores were not joined together if their measured ac capacitances differed by more than 0.3 per cent. When such capacitance differences did exist, the core length involved was removed from its normal sequence and placed in a position near the middle of the repeater section.

Because the taping and armoring processes were combined in one

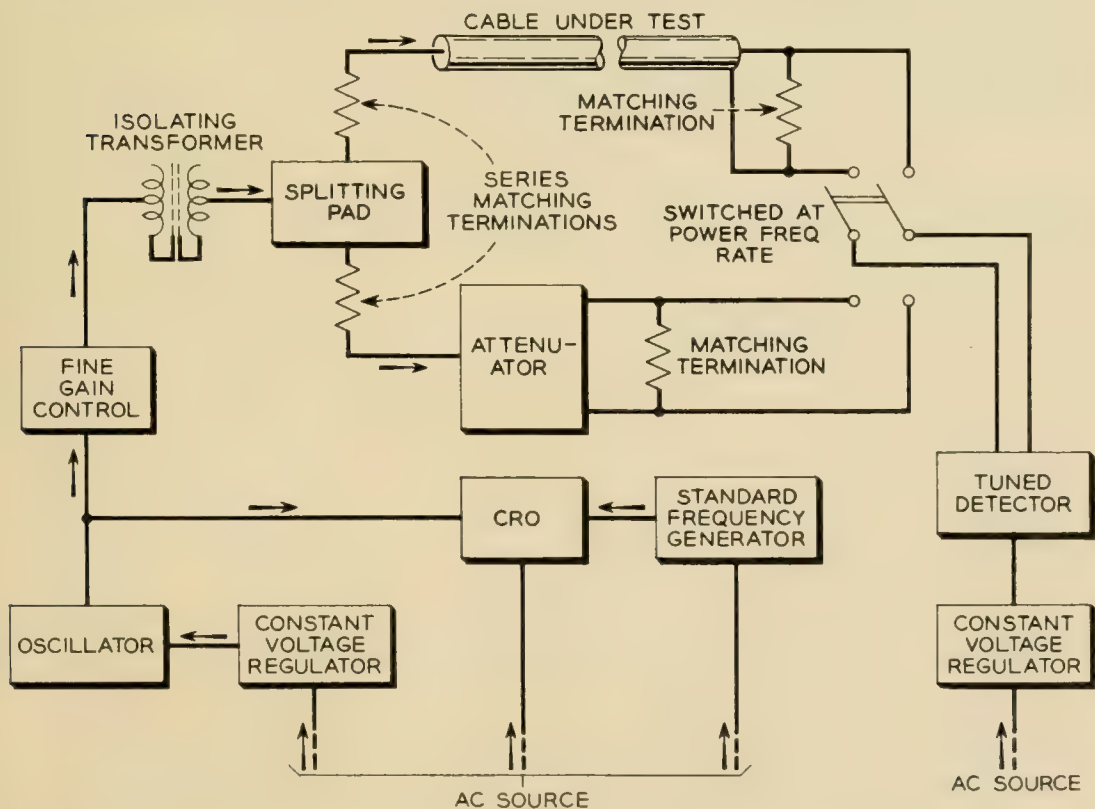


Fig. 10 — Simplified block diagram of cable attenuation measuring set.

production line in the American factory, no other electrical measurements could be made on the components of the cable until it was completely armored. In the British factory, tests for information purposes only were made on the cable in the coaxial stage. These tests included measurements of attenuation, internal impedance irregularities, and terminal impedances. They served as a means of evaluating the changes in the electrical performance during armoring.

The insulation resistance requirement after armoring and storage under water for at least 24 hours was 100,000 megohm-miles. The cable had to withstand 50,000 volts for a period of one minute without failure.

ATTENUATION MEASUREMENTS

As an aid in achieving the desired uniformity of product, new measuring equipment of improved accuracy was provided. A block schematic of this equipment is shown in Fig. 10. The requirements for this equipment were that it should be capable of measuring a 37 to 44 mile section of cable with an absolute accuracy of 0.04 db and a precision of 0.01 db in the frequency range from 1 to 250 kc.

The attenuation of the cable was measured at 10 kc intervals from 10 kc to 210 kc and measured values were corrected to 37°F, using the changes in attenuation owing to temperature, shown in Fig. 11. By comparing the corrected values with the design characteristic shown in Fig. 12, the deviations were determined. Both the attenuation charac-

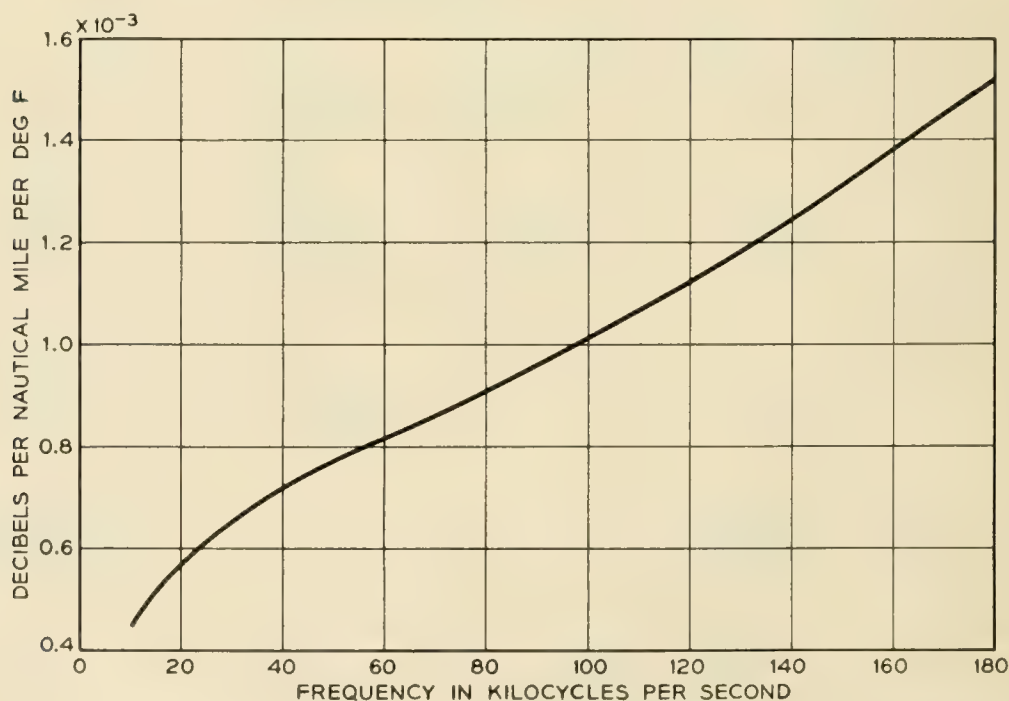


Fig. 11 — Change in cable attenuation due to temperature as a function of frequency.

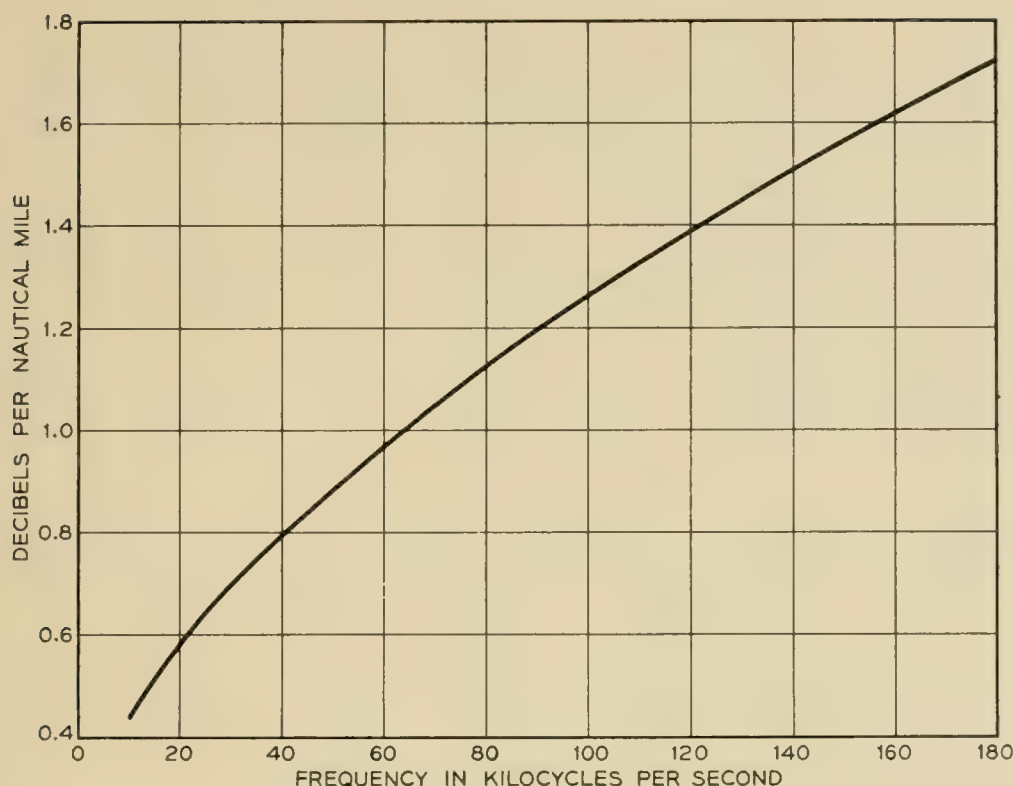


Fig. 12 — Design characteristic of cable attenuation as a function of frequency for 37°F and atmospheric pressure.

teristic and the changes in attenuation owing to temperature were derived from factory measurements of attenuation made on the Florida-Puerto Rico cable.

By comparing the running average and spread of these deviations with the design requirements, it was possible to assess the performance of the cable and, if required, to make any necessary adjustments in parameters for subsequent sections. In addition, these deviations were used to determine the length adjustment required for each repeater section to keep the sum of the deviations at each frequency in any one ocean block to a minimum. Typical average attenuation deviation characteristics are shown in Fig. 13.

TEMPERATURE AND PRESSURE COEFFICIENTS

Measurements of primary constants were made on 20-foot lengths of cable and core placed in a temperature and pressure controlled tank. These measurements were used to compute the temperature coefficients of attenuation in order to check the values derived from measurements made on the Florida-Puerto Rico cable section. Additional attenuation measurements were made on several repeater section lengths of cable over a range of temperature from approximately 40° to 70°F, to establish further the magnitude of the changes in attenuation with tempera-

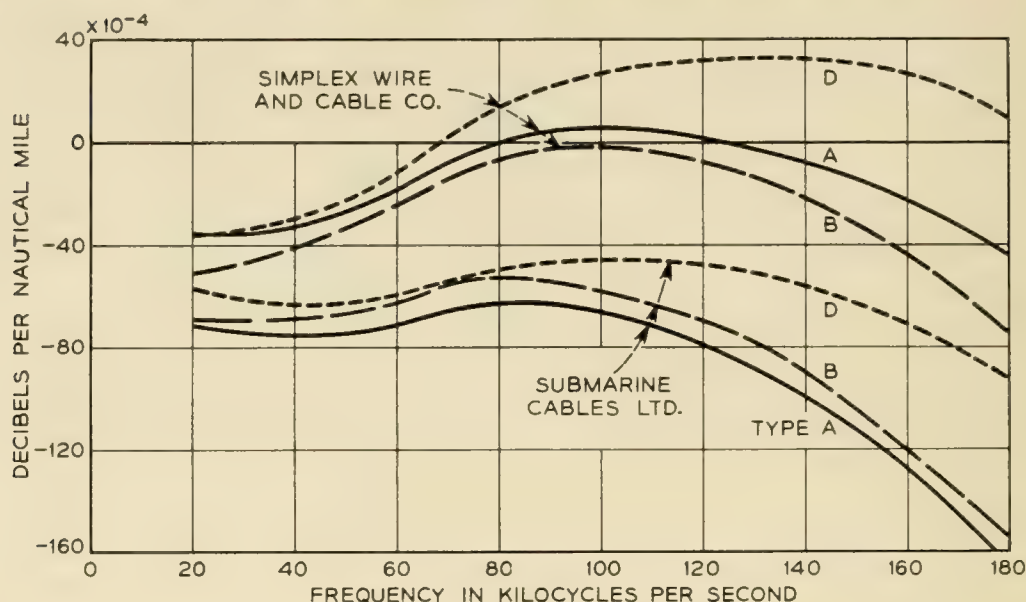


Fig. 13 — Deviation of measured cable attenuation from design characteristic as a function of frequency. Typical average values for Types A, B, and D for 37°F and atmospheric pressure.

ture. The measurements indicated that the derived temperature coefficients were accurate to within ± 10 percent.

Measurements also were made to determine the effect of pressure on the primary constants of the cable. These measurements indicated that capacitance was the only parameter affected by pressure. The capacitance increased linearly 0.1 percent for each 500 pounds per square inch of applied pressure. Since the attenuation, α , is inversely proportional to the impedance, it is evident that if C is the only parameter affected by pressure, α will also be affected by pressure to an amount equal to approximately one half the pressure effect on C . The pressure coefficient of α was therefore established as 0.05 percent per 500 pounds per square inch of pressure.

LAVING EFFECT

Analysis of the Florida-Puerto Rico cable data indicated that the measured ocean bottom attenuation was less than the attenuation predicted from factory measurements. The differences were large enough to warrant study and indicated that the measurements were in doubt or sea bottom conditions were not known accurately or that some unexplained phenomenon was taking place.

In March of 1955, approximately 22 miles of cable of the transatlantic design were laid in 300 fathoms of water off the coast of Spain in the Bay of Cadiz, and another equivalent length was laid in 2,300 fathoms off Casablanca. Precise measurements of attenuation were made in both cases, and it was established that a difference did in fact exist

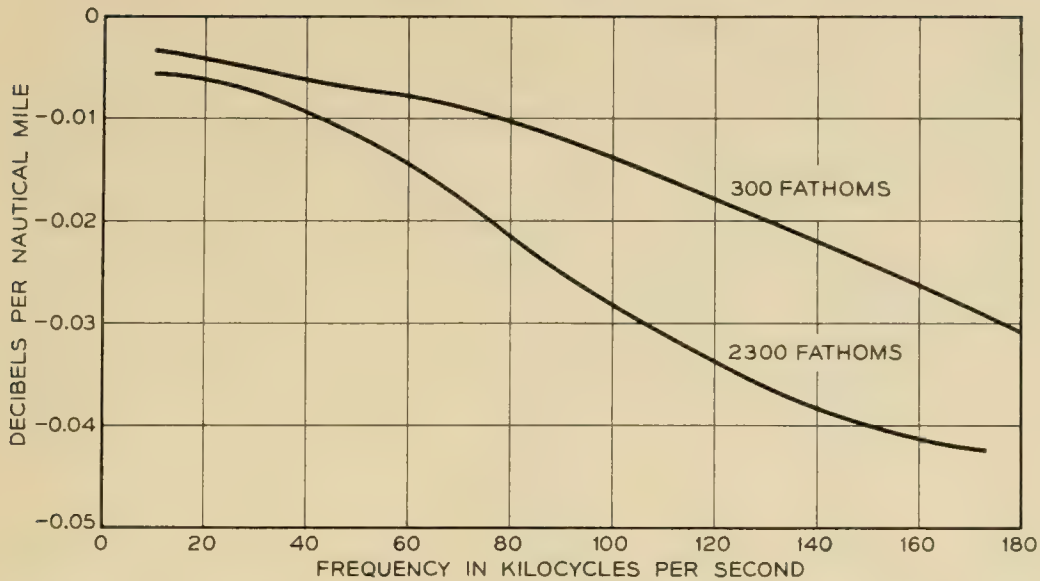


Fig. 14 — Laying effect or deviation of measured attenuation from predicted attenuation as a function of frequency, as observed in Gibraltar trials.

between measured values of attenuation at the ocean bottom and values predicted from factory measurements. It was further established that the difference in 2,300 fathoms was about twice that in 300 fathoms. The measured differences are shown in Fig. 14.

It was established during these trials that the difference increased slightly with time. Measurements made on the cable in 300 fathoms immediately, 18 hours, 48 hours, and 86 hours after laying indicated that measurable changes in attenuation were taking place. However, the change between 48 and 86 hours was so small that it was concluded only very small changes would occur in a moderate interval of time. The tests also indicated that the attenuation of the two lengths of cable decreased somewhat during loading of the cable ship.

The total difference between the attenuation at the ocean bottom and the values predicted from factory measurements, taking the temperature and pressure coefficients into account, was designated "laying effect". Various theories, such as the consolidation of the central conductor, consolidation of return structure, and changes in the dielectric material have been advanced to explain these differences. Each of these has been under study, but at the time of writing this paper, no conclusive explanation has been established.

The shape of the "laying effect" versus frequency characteristic was such that the adjustment of repeater section lengths in conjunction with several fixed equalizers, which had approximately 4 db loss at 160 kc and 0.6 db loss at 100 kc, would provide a good system characteristic. The matter of equalization is covered in greater detail in the article⁴ on the overall system. The magnitude of the laying effect observed during the laying of the two transatlantic cables substantiated the trial results.

PULSE ECHO MEASUREMENTS

Process controls, such as the use of a capacitance monitor and the jointing of the core in manufacturing sequence provided the means for controlling the magnitude of reflections due to impedance mis-matches. However, to insure that the final product met these requirements, measurements of terminal impedance and internal irregularities were made using pulse equipment.

A block schematic of the circuit of the echo set is shown in Fig. 15. For the submarine cable tests, a 1.5-microsecond raised cosine pulse

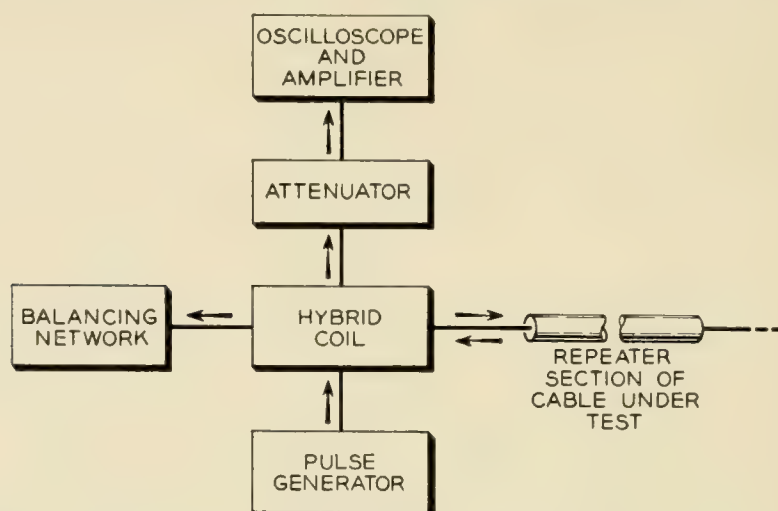


Fig. 15 — Simplified block diagram of pulse echo set for measurement of terminal impedance and internal impedance irregularities.

was used, and the impedance of the balancing network was calibrated at 165 kc. The 165-kc impedance of the repeater sections was maintained well within a range of 54.8 ± 1 ohm. The internal irregularities at the point of irregularity were maintained at least 50 db below the magnitude of the measuring pulse. The requirement was 45 db.

ACKNOWLEDGMENTS

The authors wish to acknowledge the many contributions made by the members of the staff of the Simplex Wire and Cable Company, Submarine Cables Limited, the British Post Office, and the Bell Laboratories groups involved during cable design and manufacture.

REFERENCES

1. E. T. Mottram, R. J. Halsey, J. W. Emling, and R. G. Griffith, Transatlantic Telephone Cable System — Planning and Over-All Performance. See page 7 of this issue.
2. J. J. Gilbert, A Submarine Telephone Cable with Submerged Repeaters, B.S.T.J., **30**, Jan, 1951.
3. P. T. Haury and L. M. Ilgenfritz, Air Force Submarine Cable System, Bell Lab. Record, Sept., 1956.
4. H. A. Lewis, R. S. Tucker, G. H. Lovell and J. M. Fraser, System Design for the North Atlantic Link. See page 29 of this issue.

System Design for the Newfoundland–Nova Scotia Link

By R. J. HALSEY* and J. F. BAMPTON*

(Manuscript received September 14, 1956)

The design and engineering of the section of the transatlantic cable system between Newfoundland and Nova Scotia were the responsibility of the British Post Office. The transmission objectives for this link having been agreed in relation to the overall objectives, the paper shows how these were translated into system and equipment design and demonstrates how the objectives were realized.

INTRODUCTION

Under the terms of the Agreement,¹ it was the responsibility of the British Post Office to design and engineer the section of the transatlantic cable system between Newfoundland and Nova Scotia. In common with other parts of the system, all specifications were to be agreed between the Post Office and the American Telephone and Telegraph Company, but as both the British and the American types of submerged repeater had been carefully studied and generally approved by the other party prior to the agreement, the basic pattern of the system was clear from the beginning.

The service and transmission objectives for the overall connections London–New York and London–Montreal were agreed² in early joint technical discussions in New York and Montreal and the agreed total impairments were divided appropriately between the various sections. In this way, the transmission objectives for the Newfoundland–Nova Scotia link were established.

ROUTE

The choice of Clarenville as the junction point of the two submarine sections of the transatlantic system was determined primarily in relation

* British Post Office.

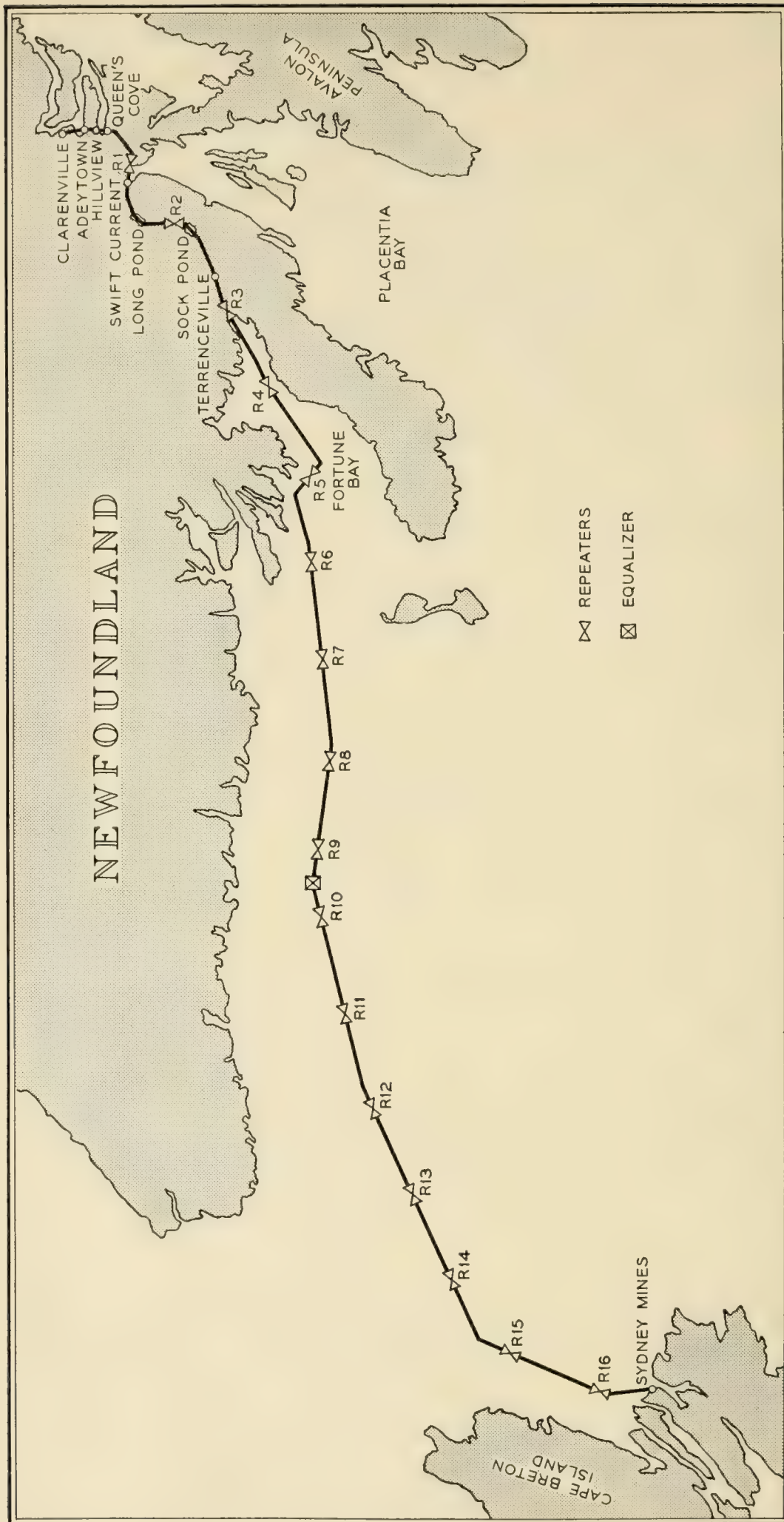


Fig. 1 — Map of route.

to the Atlantic crossing and the desire to follow a transatlantic route to the north of existing telegraph cables.³ There were a number of possibilities for the route between Clarenville and the east coast of Cape Breton Island, the most easterly point which could be reached reliably by the radio-relay system through the Maritime Provinces of Canada. One possibility, which had been considered earlier, was to cross Newfoundland by a radio-relay system and to employ a submarine-cable link across Cabot Strait only. The final decision to build a cable system between Clarenville and Sydney Mines raised a number of problems in respect of the route to be followed, concerned primarily with potential hazards to the cable brought about by:

(a) The existence of very extensive trawling grounds on the Newfoundland Banks.

(b) The location of considerable numbers of telegraph cables in the vicinity.

(c) Grounding icebergs.

The route finally selected after thorough on-the-spot investigations^{3, 4, 5} (Fig. 1) is satisfactory in respect of all these hazards, involving no cable crossings and being inshore of the main fishing grounds. The straight-line diagram of the route is shown in Fig. 2; the total cable length is 326 nautical miles, of which 54.8 nautical miles are between Clarenville and Terrenceville, Newfoundland, where the cable finally enters the sea. The maximum depth of water involved is about 260 fathoms.

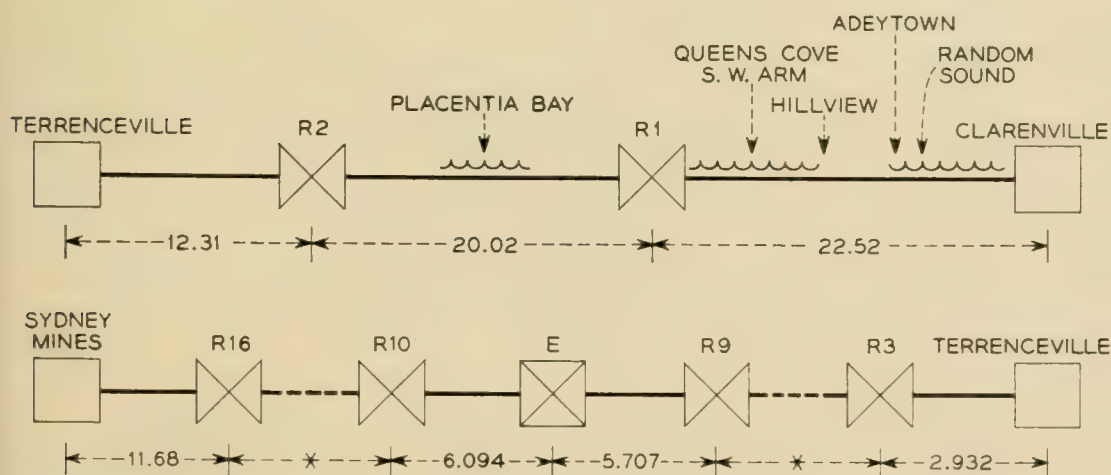


Fig. 2 — Straight-line diagram of route. R. Repeater. E. Equalizer. All distances are in nautical miles.

* Repeater spacing R3-R9 and R10-R16, 20.4 n.m.

CABLE

Choice of Design

Since 1930, when the Key West-Havana No. 4 cable was constructed,⁶ it has been usual to extrude the insulation of coaxial submarine cables to a diameter of about 0.62 inch, and most of the cables in the waters around the British Isles are of this size. The experience of the British Post Office with submerged repeaters⁷ in its home waters, dating from 1944, when the first repeater was laid between Anglesey and the Isle of Man,⁸ has therefore been mainly with 0.62 inch cables, first with paraggutta as a dielectric and later with polyethylene. Most of these cables were originally operated without repeaters, and the 60-circuit both-way repeaters which are now installed on the routes were designed to match their characteristics.

In planning a new system, the size of cable will be determined by one of the following considerations:

- (i) Minimum annual charges for the desired number of circuits.
- (ii) Terminal voltage required to feed the requisite number of repeaters.
- (iii) Maximum number of repeaters or minimum repeater spacing which is considered permissible.
- (iv) Maximum (or minimum) size of cable which can be safely handled by the laying gear in the cable ship.

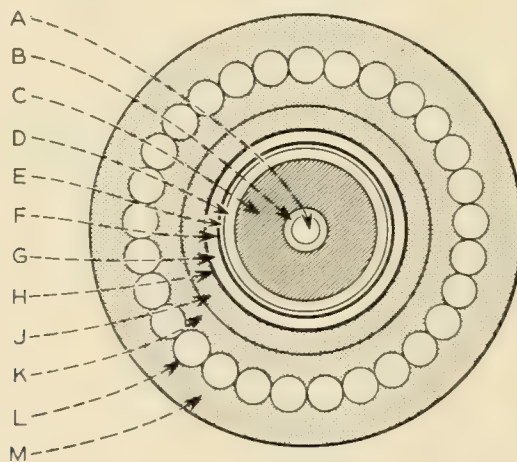


Fig. 3 — Cross-section of cable across Newfoundland showing make-up. A. Centre conductor, 0.1318-inch in diameter copper. B. Three 0.0145-inch copper surround tapes. C. Polyethylene to 0.620-inch diameter. D. Six 0.016-inch copper return tapes. E. 0.003-inch overlapped copper teredo tape. F. Impregnated cotton tape. G. Five iron screen tapes. H. Impregnated cotton tape overlapped. J. Polyethylene sheath to 1.02-inch diameter. K. Inner serving of tanned jute yarn. L. Armour wire 29 x 0.128-inch diameter. M. Outer serving of tarred jute yarn.

When the 36-circuit system between Aberdeen, Scotland, and Bergen, Norway, was planned in 1952, the route length (300 nautical miles) greatly exceeded that of any other submarine telephone system, and it was decided to use a core diameter of 0.935 inch, first, to keep the number of repeaters as low as seven, and second, because the system was intended as a prototype of a possible Atlantic cable. The cable dielectric is polyethylene (Grade 2) with 5 per cent polyisobutylene.

For the Clarenville-Sydney Mines link it proved possible to design for minimum annual charges. With increasing experience and confidence in submerged repeaters, it was no longer considered necessary to restrict the number of repeaters as for Aberdeen-Bergen, and the terminal voltage requirements were reasonable. At the current prices of cable and repeaters in Great Britain the optimum core diameter for 60 both-way circuits is about 0.55 inch, but the increased charge incurred by using 0.62-inch cable is less than 5 per cent (0.62-inch core is optimum for 120 both-way circuits). In order to facilitate manufacture and the provision of spare cable, it was therefore logical to adopt the same design as that proposed for the Atlantic crossing and described elsewhere.^{3, 9}

After investigating various possible types of cable for the overland section in Newfoundland, it was decided to use a design essentially the same as the main cable but with additional screening against external interference.⁴ As far as the outer conductor and its copper binding tape, the construction (Fig. 3) is identical with that of the main cable except that the compounded cotton tape is overlapped. Outside this are five layers of soft-iron tapes each 0.006-inch thick, the innermost being longitudinal and the others having alternate right- and left-hand lays at 45° to the axis of the cable. After another layer of compounded cotton tape there is extruded a polyethylene sheath 0.080 inch thick, and the whole is jute served and wire armoured. As a check on the efficiency of the screening, the maximum sheath-transfer impedance at 20 and 100 kc was specified as 0.005 ohm per 1,000 yards.

It was thus possible to treat the entire link from Clarenville to Sydney Mines as a uniform whole, using the same type of repeater on land as in the sea. A small hut at Terrenceville contains passive networks only.

Attenuation Characteristics

When the system was designed, precision measurements of cable attenuation were not available. The design of the Oban-Clarenville link was based on laboratory measurements on earlier 0.62-inch cable of a

similar type, but the available data applied only to frequencies up to about 180 kc, whereas the Clarendville-Sydney Mines link was to operate at frequencies up to 552 kc; extensive extrapolation was therefore involved. As soon as the first production lengths of cable became available in February, 1955, laying trials were carried out off Gibraltar, and it was found that there were serious changes of attenuation on laying, over and above those directly attributable to temperature and pressure effects, and that the assumed characteristics were inaccurate. Although the attenuation in the factory tanks had been in reasonable agreement with that of the earlier cable, there were changes on transfer to the ship's tanks and again on laying, amounting in all to a reduction of about 1.5 per cent at 180 kc. This would have been comparatively unimportant had the discrepancy been of 'cable shape', i.e., the same fraction of the cable attenuation at all frequencies and therefore exactly compensated by a length adjustment of the repeater sections. As this was not so, and as the cable-equalizing networks in the repeaters were settled by this time, it was clear that precise information must be obtained in order that suitable additional equalizers could be provided for insertion in the cable on laying. There are a number of factors which can lead to small changes of attenuation on laying, but most of these tend to increase the losses. The primary reason for the observed changes appears to be contact variations between the various elements of the inner and outer conductors, i.e. the wire and three helical tapes forming the centre conductor and the six helical tapes forming the return conductor. These contact resistances tend to change with handling, and as a result of a slight degree of 'bird caging' when coiled, it seems that the attenuation decreases as the coiling radius increases, and vice versa. Also, the effect of sea pressure is to consolidate the conductors and thereby further reduce the attenuation — an effect which appears to continue on a diminishing basis for a long time after laying.

To obtain reliable data for the Clarendville-Sydney Mines link, 10 nautical miles of cable with A-type armour was laid at about the mean depth of the system (120 fathoms), off the Isle of Skye. The attenuations, coiled and laid, are shown in Fig. 4, due allowance having been made for temperature and pressure. The ordinates — attenuation versus frequency — are such that the value should be approximately constant at high frequencies.

In making a final determination of the cutting lengths for the repeater sections, it was assumed that the factory measurements of attenuation would be reduced by 1.42 per cent at 552 kc on laying, that the temperature coefficient of attenuation would be +0.16 per cent

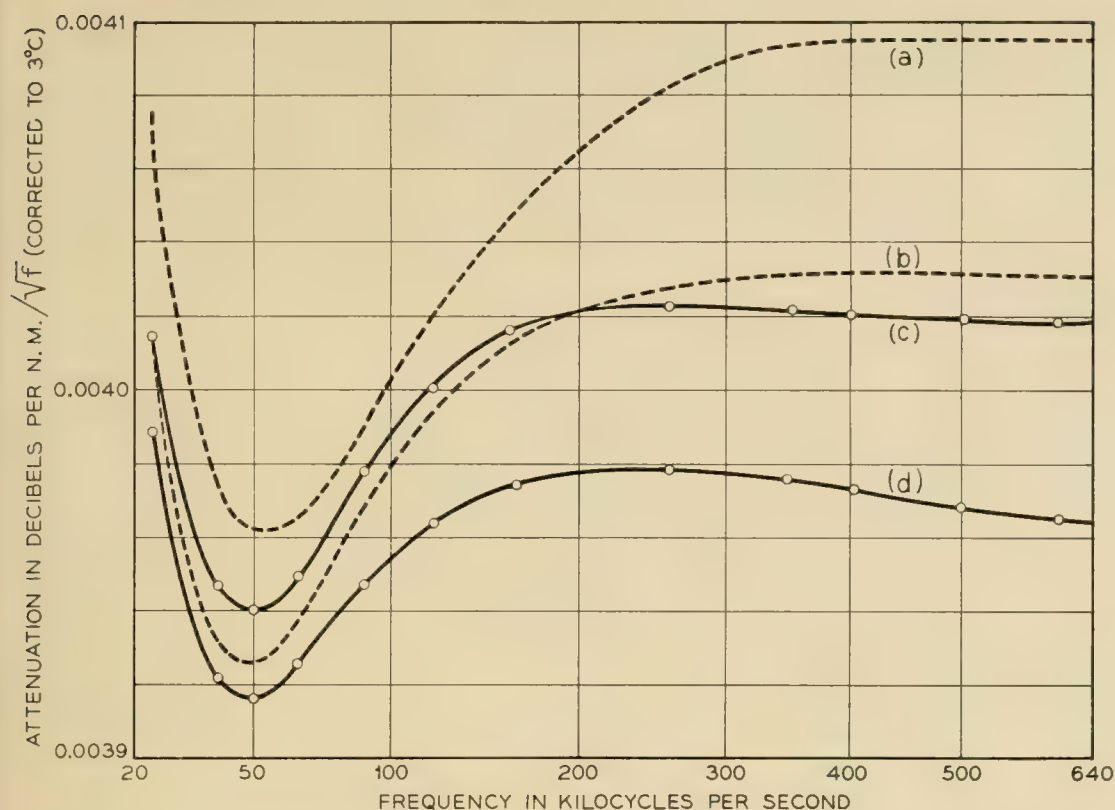


Fig. 4 — Cable attenuation characteristics — Skye trials. (a) Characteristic originally assumed. (b) Characteristic measured in factory tank (flooded). (c) Characteristic measured in ship's tank (flooded). (d) Characteristic measured after laying.

per deg C and that the true pressure coefficient of attenuation was negligible at the depths involved.

DESCRIPTION OF SYSTEM

Circuit Provision and Frequency Allocation

It was originally thought that a design similar to that of the Aberdeen-Bergen system would be suitable for the Clarenville-Sydney Mines route in that it would provide more circuits (36) than the long section across the Atlantic. This potential excess capacity, which was required for circuits between Newfoundland and the Canadian mainland, disappeared when it was found that 36 circuits could, in fact, be provided over the longer link. The Aberdeen-Bergen design was therefore modified to provide a complete supergroup of 60 circuits, the same capacity as the earlier British projects.⁷ The system thus requires broad-band transmission of 240 kc in each direction.

In the earlier British projects the frequency bands transmitted are 24–264 and 312–552 kc, but for the present purpose the lower band is

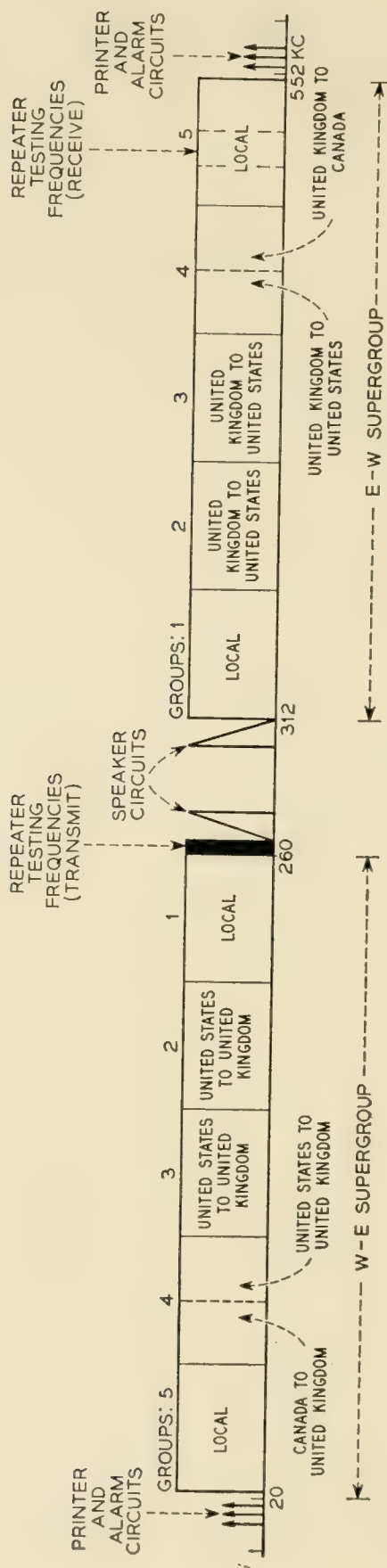


Fig. 5 — Frequency allocation.

dropped by 4 kc to 20–260 kc, so that the lowest frequency is the same as on the Atlantic cables. This enables common frequency-generating equipment to be used at Clarenville for the two links and minimizes crosstalk problems. The main transmission bands and the allocation of the five 12-circuit groups are shown in Fig. 5, together with the ancillary channels; the facilities provided are discussed later.

Submerged Repeaters

The submerged repeaters employed are fully described elsewhere,¹⁰ and it will suffice to note here that they are rigid units, approximately cylindrical in shape, 9 feet long and 10½ inch maximum diameter. They are capable of withstanding the full laying pressure in deep water, although this of little importance in the present application.

They are arranged for both-way transmission through a common amplifier which has two forward paths in parallel, with a single feedback path. The two halves of the amplifier are so arranged that practically any component can fail in one, without affecting the other.

Power-Feeding Arrangements

The submerged repeaters are energized by constant-current dc supplies between the center conductor and ground, the power units at the two ends being in series aiding and the repeater power circuits being in series with the center conductor, i.e., without earth connections, as in Fig. 6. This is the only arrangement by which it is possible to control the supply accurately at every repeater, the insulation resistance of cable and repeaters being sufficiently great that the current in the center conductor is virtually the same at all points. The constant-current feature of the supply ensures that repeaters cannot be overrun in the event of an earth fault on the system.

On the Oban-Clarenville link the anode voltage is derived from the drop across the electron tube heaters. This results in the heaters being at a positive potential with respect to the cathodes, a condition which tends to break down the heater-cathode insulation.¹¹ In the American electron tubes this insulation is very robust and the risk is considered to be negligible, but in the current British electron tubes, which have a much higher performance, the arrangement is undesirable. In view of the much smaller number of repeaters it was possible to derive the heater and anode supplies as in Fig. 6 and thus to reverse the sense of the heater-cathode voltage and also to provide an anode voltage of 90, against 55 in the longer link.

With this arrangement the link requires a total supply voltage of about 2,300. The power-feeding equipment¹² at each terminal station is designed to feed a constant current of 316 ma at this voltage, and it is permissible to energize the system from one end only, if necessary. The repeater capacitors — the limiting factors in respect of line voltage — are rated very conservatively at 2,500 volts, so that a single-ended supply of 2,300 volts, with the possibility of superimposed ground-potential differences, is near the desirable maximum. The two terminal power units are therefore designed to operate in series and to share the voltage.

With access to the cable provided at Terrenceville it is possible to operate the power system on the following bases:

(a) No ground at Terrenceville, power from both ends on a master-and-slave basis (to ensure that the constant-current units do not build up an excessive voltage); this is the normal arrangement.

(b) No ground at Terrenceville, power from one end only.

(c) Ground at Terrenceville, with the Clarenville and Sydney Mines

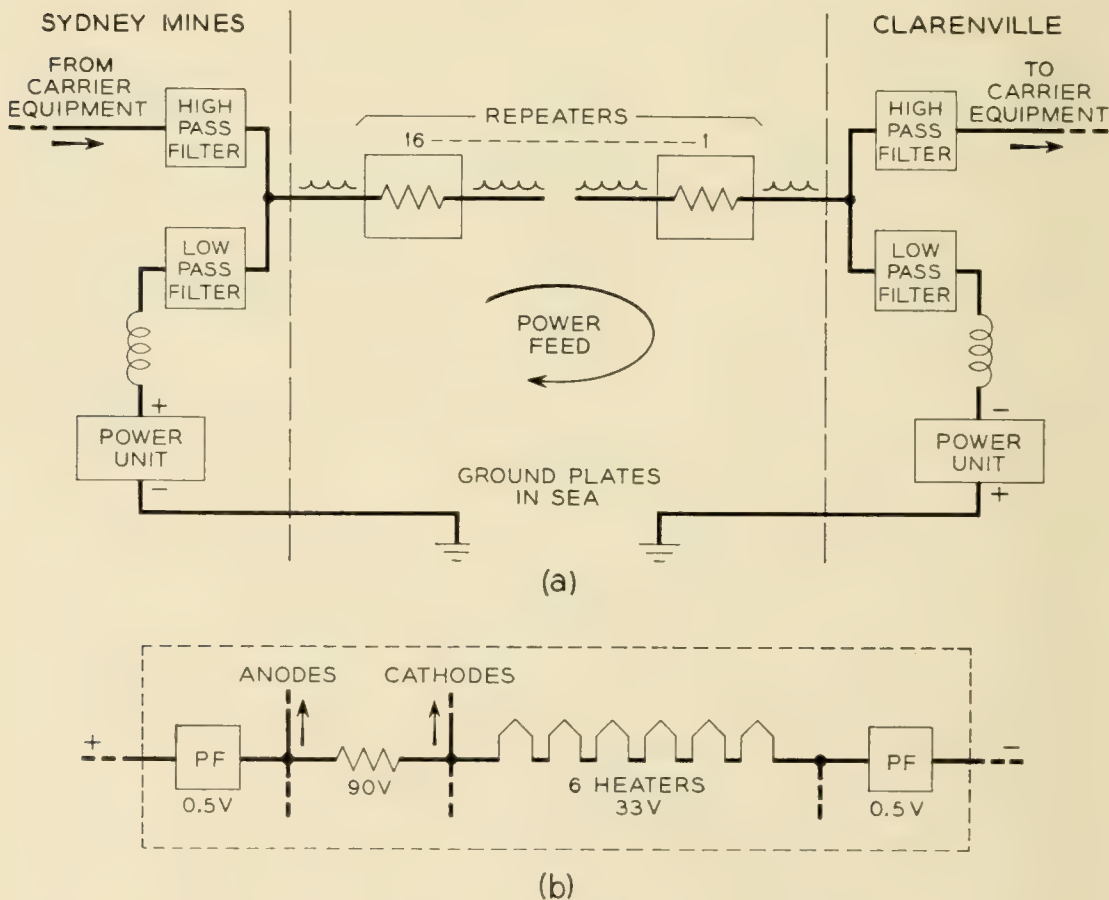


Fig. 6 — Power-feeding arrangements. (a) General schematic. (b) Repeater power circuit. P.F. — Power filter.

power units energizing the land and sea cables respectively; this arrangement has been particularly useful during the installation period.

The presence of high voltages on the cable constitutes a potential danger to personnel, hence special precautions are taken in the design of the equipment in which the cable terminates and in which high voltages exist or may exist.

The ground connections for the power circuits at the two ends are via special ground cables and ground plates located about half a mile from the main cable, and metering arrangements are provided to check that the current does in fact take this path. These measures ensure that the current returning via the cable armour is never sufficient to cause serious corrosion.

Arrangement of Terminal Equipment

Fig. 7 show the arrangement of the terminal equipment. In accordance with an early agreement defining precisely the various sections of the project, the link is considered to terminate at the group distribution frames at Clarenville and Sydney Mines, i.e. at the 60–108 kc interconnection points.

In addition to the cable-terminating and power-feeding equipments (A and B), the following are provided at the terminals:

(a) Submarine-cable terminal equipment (C) consisting of repeaters to amplify the signals transmitted to and received from the cable, equalizers and frequency-translating equipment to convert the line frequencies to basic supergroup frequencies (312–552 kc).

(b) Group-translating (group-bank) equipment (D) to convert the basic supergroup to five separate basic groups (60–108 kc) and vice versa.

(c) Equipment for the location of cable and repeater faults (E and F).

(d) Speaker and printer circuit equipments (G and H) to provide two reduced-bandwidth telephone circuits, two telegraph circuits and one alarm circuit for maintenance purposes. It is clear that such circuits should be substantially independent of the main transmission equipment.

Two principles were agreed very early in the planning; first, that the engineering of the various links should be integrated as far as possible, and second, that the items of equipment at each station should be provided by the party best in the position to do so. In consequence:

(i) Items of standard equipment were provided by the A.T. and T. Co. at Clarenville and Sydney Mines (and by the Post Office at Oban), thus simplifying maintenance and repair problems.

(ii) Basic power plant and the carrier supplies for supergroup and group translation were provided by the A.T. and T. Co. for both terminal equipments at Clarenville.

(iii) Terminal equipment special to the Clarenville-Sydney Mines link was provided by the Post Office.

In view of the importance of the link, the power-feeding and transmission equipment are completely duplicated.

The submarine-cable terminal equipment is arranged to transmit the basic supergroup, directly over the cable in the east-to-west direction. In the west-to-east direction the supergroup is translated to the range 20–260 kc, using a 572 kc carrier.

DESIGN OF TRANSMISSION SYSTEM

Performance Requirements

The agreed transmission objects for the Clarenville-Sydney Mines link were as follows:

Variation of Transmission Loss.

The variation in the transmission loss of each group should have a standard deviation not greater than 0.5 db; this implies that the variation from nominal should not exceed 1.3 db for more than 1 per cent of the time.

Attenuation/Frequency Characteristics.

Only the overall characteristics of the individual circuits were precisely specified, the limits being the C.C.I.F. limits for a 2,500-km circuit and the target one-half of this. With this objective in view, the group characteristics in each link must clearly be as uniform as is practicable.

Circuit Noise.

The total noise contributed by the link to each channel in the busy hour (i.e., including intermodulation noise) should have an r.m.s. value not exceeding +28 dba* (corresponding to –56 dbm) at a point of zero relative level.

* This refers to the reading on a Bell System 2B noise meter (FIA weighting network); the noise level (dba) is relative to a 1 kc tone at –85 dbm. In Europe, noise is measured on a C.C.I.F. Psophometer (1951 weighting network), which is calibrated in millivolts across 600 ohms; this is commonly converted to picowatts (pw). The white noise equivalence of the two instruments is given by $\text{dba} = 10 \log_{10} \text{pw} - 6 = \text{dbm} + 84$; the agreed limit of +28 dba is therefore equivalent to 2513 pw (1.23 mv or –56 dbm). The corresponding C.C.I.F. requirement at 4.0 pw/km would be 2,400 pw, this value not to be exceeded for more than 1 per cent of the time.

Crosstalk.

The minimum equal-level crosstalk attenuation should be 61 db for all sources of potentially intelligible crosstalk; this was accepted as a target for both near- and distant-end crosstalk. Although go-to-return crosstalk is not important for telephony (it appears as sidetone) and a limit of 40 db is satisfactory even for voice-frequency telegraphy, it assumes great importance for both-way music transmission; also, it was desired to be non-restrictive of future usage.

Assessment of Requirements

The design of the high-frequency path to meet the agreed requirements involves consideration of:

(a) Noise, including fluctuation (resistance and tube) noise and intermodulation.

(b) Wide-band frequency characteristics, including the effects of the directional filters at the terminal and in the repeaters.

(c) Variations of (a) and (b) in respect of temperature and aging.

The noise requirement is by far the most important factor in the design of the line system.

The choice of route and cable having been made, the total loss was known and it was necessary to determine the minimum number of repeaters to compensate for this loss and to meet the noise requirement with adequate margin for inaccurate estimates of cable attenuation after laying, temperature variations, aging and repairs. An attempt to achieve the necessary gain with too few repeaters would result in excessive noise.

Design of the amplifiers in the British repeaters is such that, with both forward paths in operation, the overload point is about +24 dbm, and with a loading of 60 channels in each direction, this permits planning levels of about -4 dbm at the amplifier output after allowing reasonable margins for errors and variations.¹⁰ Previous experience shows that, at such output levels, intermodulation noise can be neglected and the full noise allowance allotted to fluctuation noise. The effect of tube noise is to increase the weighted value of resistance noise by about 1 db to -137.5 dbm, or -53.5 dba, at the input to the amplifier in each repeater.

At the highest transmitted frequency the equalizers, power filters and directional equipment introduce losses of about 1 db and 4 db at the

input and output of the amplifier respectively; these losses must, effectively, be added to the loss in the cable.

Two other pieces of information are necessary before the repeater system can be planned — the permissible transmitting and receiving levels at the shore stations. The transmitting equipment provided at Clarenville can be operated at channel levels up to +20dbm, and it is logical to allow the same receiving level at the shore end as at intermediate repeaters.

On the above basis it is possible to construct a curve (Fig. 8) relating the total circuit noise to the number of intermediate repeaters, and it is seen that the minimum number is 15, each of which must have an overall gain of 59 db (amplifier gain, 64 db) at 552 kc. The actual provision is 16 repeaters, each having a gain of 60 db at 552 kc, the additional gain being absorbed in fixed and adjustable networks at points along the route, as indicated in the following section.

Level Diagram

The actual level diagram (planning levels are shown in Fig. 9) differs somewhat from that which can be deduced directly from the preceding because of the following considerations:

- (a) The location of the first repeater from Clarenville (i.e., on land)

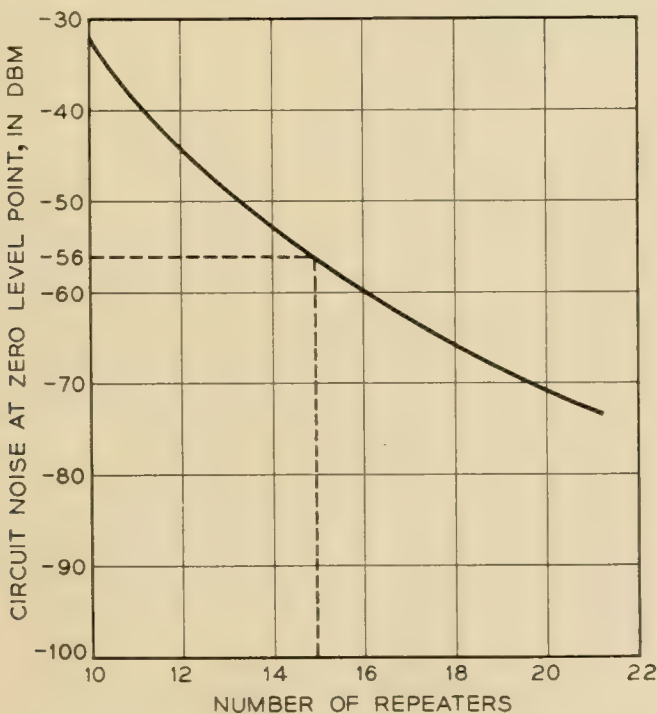


Fig. 8 — Variation of circuit noise with number of repeaters.

was dictated by topography and the desire to locate both it and the second repeater in ponds; thus the transmitting level at Clarendville is substantially lower than the permissible maximum.

(b) There are equalizing networks at Terrenceville and facilities for their adjustment to compensate for temperature variations.

(c) Because of the difference between the actual cable attenuation and that for which the repeaters were planned (see section on *Attenuation Characteristics*) it was necessary to include an equalizer unit in the sea, midway between Terrenceville and Sydney Mines. Loss equivalent to 9 nautical miles of cable was also introduced at this point to ensure that repeater No. 16 would be sufficiently far from the shore at Sydney Mines.

(d) Cable simulators are included in the cable-terminating equipment at Sydney Mines to build out this section to a standard repeater section; the actual cable length was, of course, unknown until the cable was complete.

Taking into account the existence of the intermediate networks, the repeater spacing is such that when both land and submarine cable sections are at mean temperature the compensation is as accurate as possible. In general the highest frequency is of greatest importance in this respect. Since the low-frequency channels experience less attenuation than the high-frequency channels, it is permissible to transmit them at a somewhat lower level, thereby increasing the load capacity of the amplifiers which is available to the high-frequency channels.

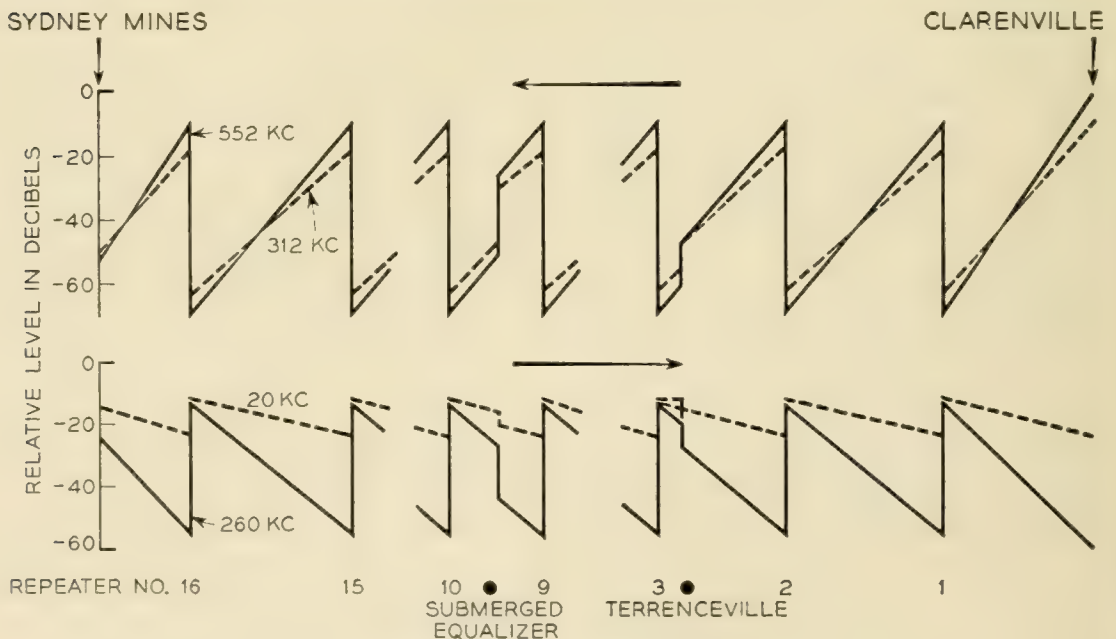


Fig. 9 — System level diagram.

Temperature Effects and their Compensation

The effect of temperature changes is likely to be somewhat complex. The land-and-sea cable sections are expected to behave in different ways in this respect, but data on the manner of variation are not very precise. The submarine cable crosses Cabot Strait, where melting icebergs drifting down from Labrador as late as June can be expected to keep the sea-bottom temperature low until well into the summer; temperatures just below 0°C were, in fact, recorded when the cable was laid in May. On land, the cable is buried 3 feet deep in bog and rock, and traverses many ponds; some data on temperatures under similar conditions in other parts of the world were available.

For planning purposes it was clear that the assumptions made would have to be somewhat pessimistic, and the assumed ranges of temperature, with the corresponding changes of attenuation at 552 kc, were:

Sea section		$2.3 \pm 3^{\circ}\text{C}$;	$\pm 4\text{ db}$
Land section		$7.5 \pm 10^{\circ}\text{C}$;	$\pm 3\text{ db}$

A possible method of circuit adjustment for temperature changes is to increase the gains equally at the sending and receiving terminals as the temperature rises and to reduce them equally as it falls. Under such conditions the effect of temperature variations on resistance noise is not very important; the levels at repeaters near the center of the route remain substantially constant, and the increase in noise from the repeaters whose operating levels are reduced is partly compensated by the reduction in noise from those whose levels are increased. The effect of the level changes on repeater loading is, however, more important as it is undesirable that any repeater in the link should overload, and additional measures which can be readily adopted to avoid serious changes in repeater levels are clearly desirable.

The estimated change of attenuation of the land sections is seen to be roughly equal to that of the submarine section, so that, from the point of view of temperature changes, Terrenceville is near the electrical center of the link. It was thus both desirable and convenient to provide adjustment at this point: Fig. 10 illustrates the advantage of seasonal changes in equalizer setting at Terrenceville, showing the way the output levels of repeaters are likely to vary along the route. The system of temperature compensation adopted therefore involves adjustable networks at both ends of the system and at Terrenceville. All the networks are cable simulators; hence the process of temperature compensation consists, effectively, in adding 'cable' when the temperature falls and removing it when the temperature rises.

At Terrenceville, the networks permit adjustments equivalent to ± 1 nautical mile of cable (3 db at 552 kc), but at Clarendville and Sydney Mines adjustments equivalent to 0.5 db at 552 kc are provided. It should therefore always be possible to maintain the overall loss of the system within ± 0.25 db, and the level at any repeater should never change by more than ± 2 db.

System Pilots

The use of pilot tones applied at constant level at the input of a system with indicating or alarm meters at the receiving end is standard on land systems on both sides of the Atlantic, although the philosophies underlying the methods of use differ. On the submarine cables round the British Isles, with or without submerged repeaters, pilot tones are used

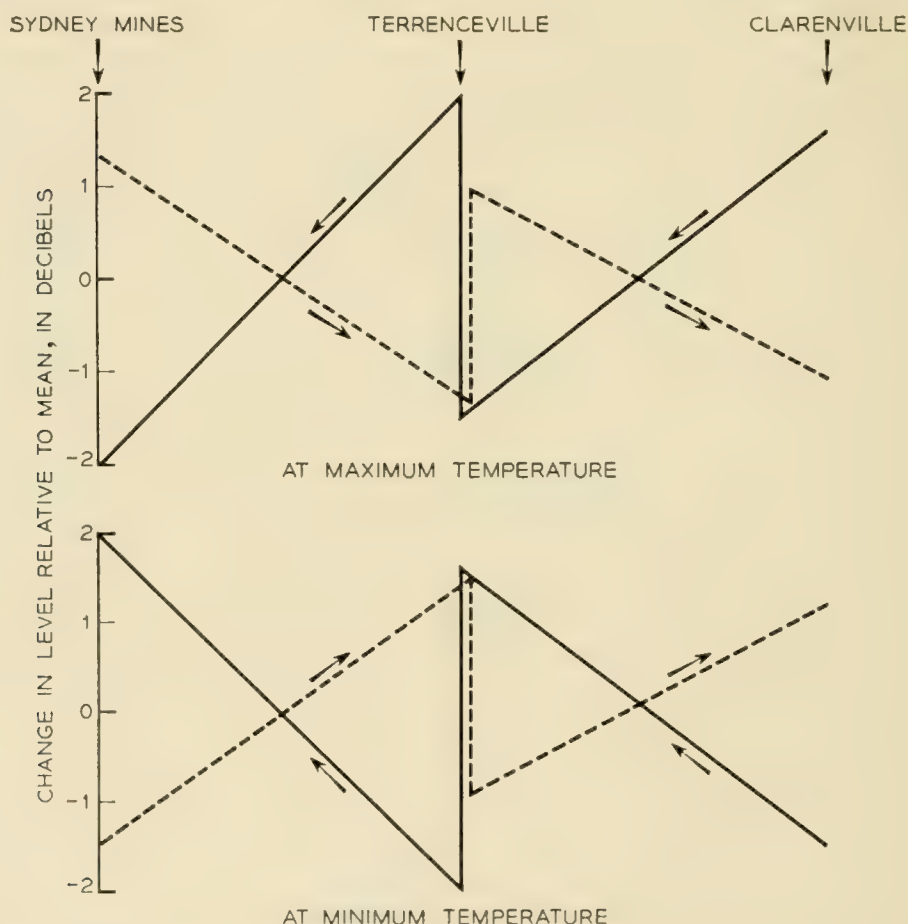


Fig. 10 — Deviation from mean of transmission levels with optimum adjustments of equalizers at Sydney Mines, Terrenceville and Clarendville.

----- W-E at 260 kc.

———— E-W at 552 kc.

Maximum deviation in the two directions occurs at the above frequencies.

to indicate the attenuation of the transmission path; these pilots are normally located just outside the main transmission bands in each direction. In the Clarenville-Sydney Mines system the frequency bands just outside the main transmission bands are occupied by telephone speaker and teleprinter circuits and by monitoring frequencies associated with the repeaters (see Fig. 5); this prevents the use of out-of-band pilots.

Fortunately, the standard Bell System group equipment is designed to apply 92-kc pilots to each group and to measure the corresponding received level. Although these are essentially group pilots, being applied and measured at points in the 60–108-kc band, it was decided that they could reasonably replace the out-of-band pilots. These pilots are blocked at each end of the system and therefore function as section pilots only.

Normal Post Office practice, both on land and submarine systems, is to use recording level meters to provide a continuous and permanent record of the pilot levels. In the present system such recording meters are used on the 92-kc pilots of two groups in each direction of transmission.

In addition to the section pilots the system carries the 84.080-kc end-to-end pilots in each of the three transatlantic groups.

MAINTENANCE FACILITIES

Speaker and Printer Circuits

It was part of the planning of the transatlantic system that two low-grade telephone (speaker) and two telegraph (printer) circuits should be provided over the submarine cables, outside the main transmission bands, and that the speaker circuits in particular should be reasonably independent of the main terminal equipment. One speaker circuit is required for local communication between the terminals of each section, the other to form part of an omnibus circuit connecting the principal stations on the route including Montreal. The arrangement for teleprinter communication was that one channel should be an overall all-station omnibus printer, the other being a direct London-New York printer.

Independent frequency-translating equipment is provided to connect the speaker and printer bands (each 4 kc) to the line. The carrier frequencies required for the speaker are provided by independent oven-controlled crystal oscillators, but for the printer the independent generation of high-stability 572-kc carriers was not considered to be justified and the main station supplies are used.

Two half-bandwidth telephone circuits are provided in the 4 kc speaker

bands by the use of standard A.T. and T. band-splitting equipment (EB banks). Signalling and telephone equipment are provided to give the required omnibus facilities on one circuit and local-calling facilities on the other. The arrangement of the speaker and printer equipment at Sydney Mines is shown in Fig. 11.

In the telegraph band a third channel transmits an alarm to the remote terminal when the 92-kc pilots incoming from that terminal fail simultaneously.

Fault Location

The speedy and accurate location of faults in repeatered cables is of very great importance, owing to the number of circuits involved and the difficulty and cost of repairs. The standard dc methods which have been applied in the past to long telegraph cables are, of course, available. The application of these methods is, however, recognized as being rather more

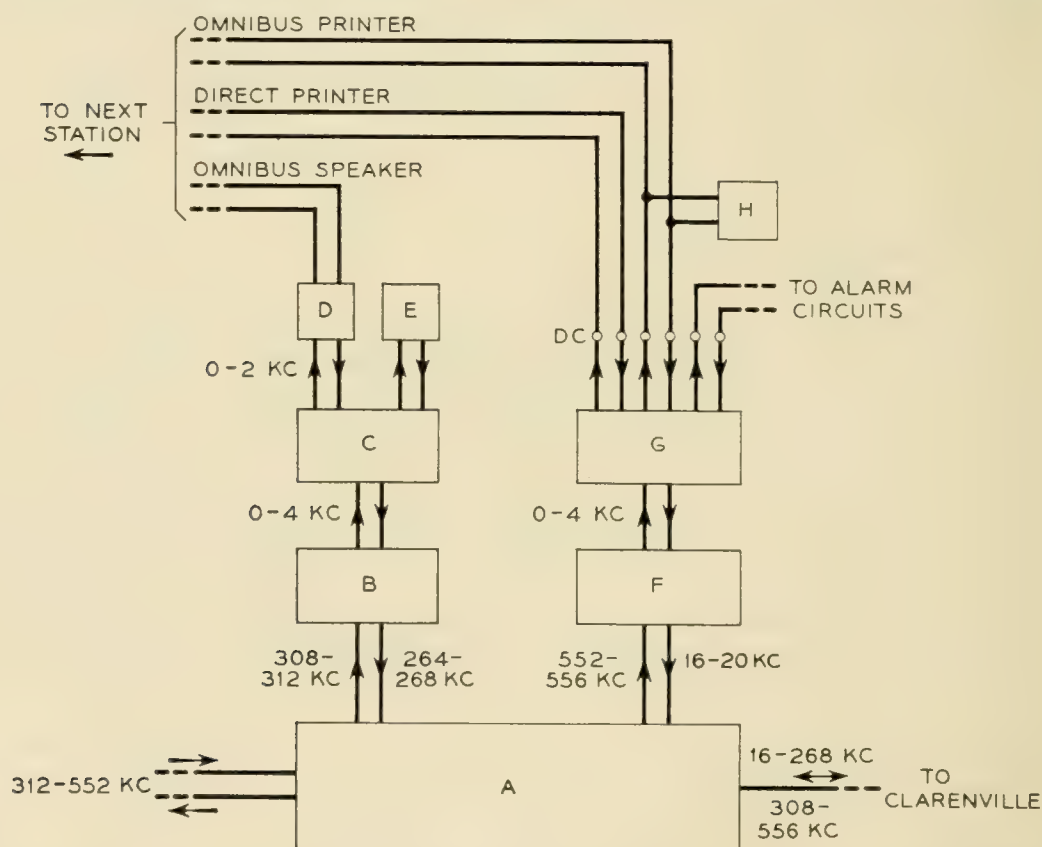


Fig. 11 — Arrangement of speaker and printer equipment at Sydney Mines. A. Submarine cable terminal equipment. B. Speaker circuit equipment. C. Emergency-band bank equipment. D. Omnibus speaker telephone. E. Local speaker telephone. F. Printer circuit equipment. G. Three-channel telegraph equipment. H. Printer.

in the nature of an art than a science and usually requires an intimate knowledge of the behaviour and peculiarities of the particular cable concerned. While the problem appears at first sight to be simple it is complicated by:

(a) The presence of ground-potential differences along the cable, sometimes amounting to hundreds of volts; these vary with time.

(b) Electrolytic e.m.f. generated when the center conductor is exposed to sea water.

(c) Absorption effects in the dielectric of the cable.

When repeaters are added the position is further complicated by:

(d) The lumped resistance of the repeaters, which is current-dependent and exceeds the cable resistance.

(e) The lumped capacitance of the repeaters with an absorption characteristic which differs from that of the cable.

It is a great advantage of both-way transmission over one cable that, by introducing some form of frequency changer at each repeater, signals outgoing in one direction can be looped back to the sending terminal. There have been a number of developments based on this principle, and in the Clarendville-Sydney Mines link two methods are available for use. Of these, the so-called 'loop-gain' method uses steady tones and depends on selective frequency measurements to discriminate between repeaters; the second is a pulse method in which repeaters are identified on the basis of loop transmission time.

The use of these methods under fault conditions depends on the possibility of keeping the repeaters energized. Work is in progress to develop methods of fault location which are of general application and do not depend on the activity of the repeaters, but these are outside the scope of the present paper.

Loop-Gain Method.

In the loop-gain method, the frequency changer in the repeater takes the form of a frequency doubler and each repeater is identified uniquely by one of a group of frequencies spaced at 120 cycles and located immediately above the lower main transmission band in the frequency range 260-264 kc. Since the frequency changing is in an upward sense, the measuring terminal is Sydney Mines, which transmits the lower band. On the Clarendville side of the directional filters in each repeater is connected, via series resistors, a crystal filter accepting the test frequency appropriate to the repeater [see Fig. 12(a)]; this frequency is doubled, filtered and returned to the repeater at the same point at which

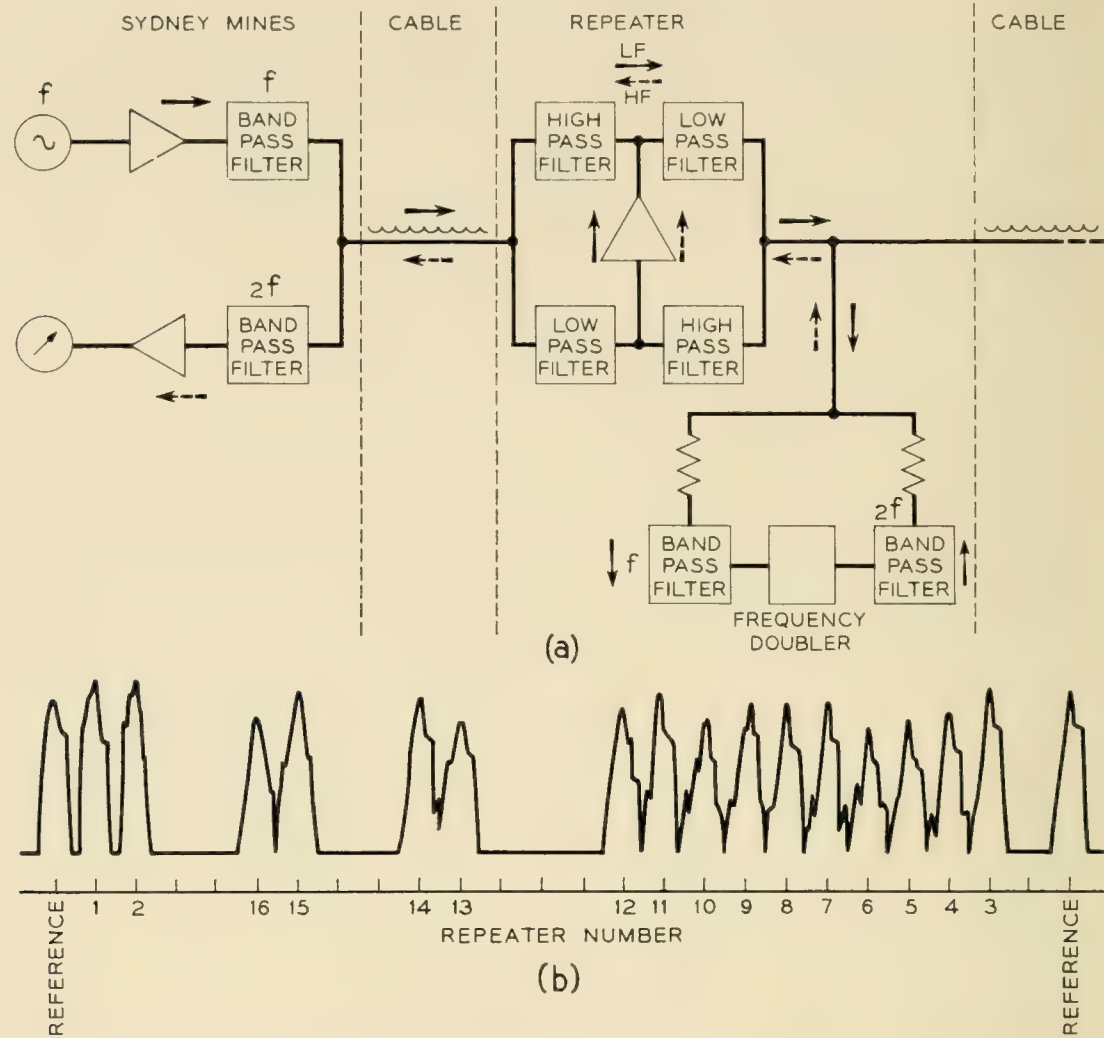


Fig. 12 — Fault location — loop-gain method. (a) Block schematic. Frequency f is in the band 260-264 kc. (b) Diagram of display.

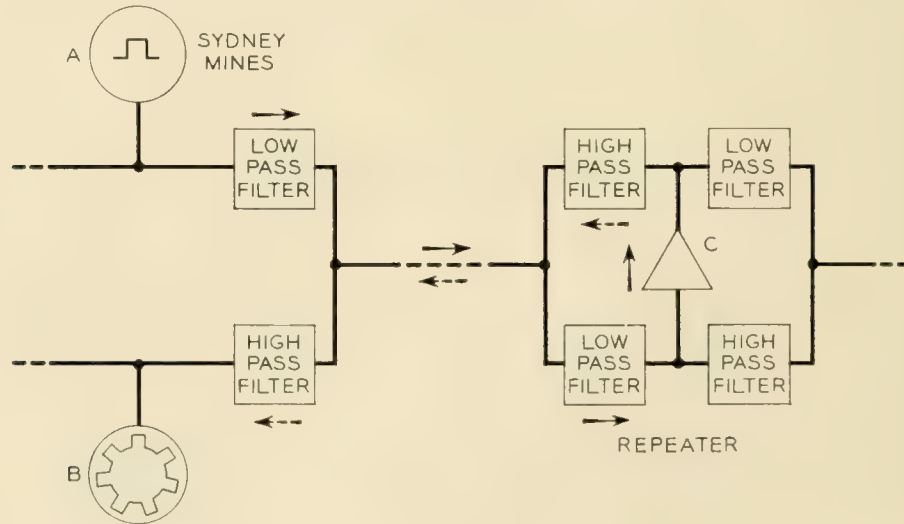


Fig. 13 — Fault location: pulse method. A. Pulse generator. B. Display of received pulses. C. Point of intermodulation (output of amplifier).

the original frequency is selected. From this point it passes via the high-pass directional filters through the amplifier and back to Sydney Mines, where the level is measured on a transmission measuring set. Information obtained in this way on each repeater can be compared at any time with that obtained when the system was installed and any gain variations localized. The test equipment provided has the additional facility of an automatic sweep of the test frequency at 4 cycles and a display of the returned-signal levels on a cathode-ray tube as in Fig. 12(b).

Although the transmitted signals lie outside the band of the W-E supergroup, the received signals, 520–528 kc, lie within the band of the E-W supergroup, and two channels must be removed from traffic to carry out the tests; these channels are in a “local” group.

Pulse Method.

As applied to the present system, the pulse method utilizes the overload characteristic of the amplifier to effect the frequency change in the repeater. At Sydney Mines a continuous train of single-frequency pulses is applied in the lower transmission band, such that either the second or third harmonic is returned in the upper band, as in Fig. 13; at Clarenville two-frequency pulses are applied in the upper band such that either a second- or third-order difference product is returned in the lower band. The pulse length is 0.15 millisecond, and the frequencies used are given in Table I. At Sydney Mines the signals can be sent and received either on the line itself or via the group equipment; in the latter case only one group need be taken out of service. At Clarenville line measurements only are provided for.

The primary display is on a cathode-ray tube with a circular time-base, and any one returned pulse can be accurately compared with the reference pulse on a second tube with a linear time-base. The pulse selected for such measurement is automatically blacked out on the primary display.

TABLE I

Station	Send to line		Product	Receive
	f_1	f_2		
	kc	kc		kc/s
Sydney Mines	216	—	$2f_1$	432
	144	—	$3f_1$	432
Clarenville	530	380	$f_1 - f_2$	150
	530	340	$2f_2 - f_1$	150

Usefulness of the Loop-Gain and Pulse Methods.

Both methods require that all the repeaters between the testing terminal and the fault can be energized. If the fault is in the cable there is a very high probability that the center conductor will be exposed to the sea, in which case the power circuit can be maintained on one side of the fault at least, although it may be somewhat noisy. Because the system is short, it is permissible to energize the link fully from one end only. The condition can never arise — as it can in the Oban-Clareville link — that the line current is limited by the maximum permissible terminal voltage.

The loop-gain test is concerned with the amplifiers in their linear regime and gives no indication of the overload point; for this the pulse test must be used. On the other hand, the pulse test does not permit accurate measurement of levels, since the pulse level reaching a particular repeater may be restricted by the overload of an earlier repeater in the chain. The pulse test is particularly useful in providing a check that both sides of each amplifier are in operation and in locating a fault of this type.

Each method depends for its operation on non-linearity at a point within each repeater and can only identify a fault as lying between two such consecutive points in the link. It is therefore desirable that these points should be as close as possible to the terminals of the repeater in order to ensure that the faulty unit can be identified. In this respect the loop-gain test has the advantage over the pulse test.

EXECUTION OF WORK

Problems due to the remoteness of the site were overcome without undue difficulty with the co-operation of the other parties concerned in the project, but the present paper would be incomplete without a brief reference to the cable- and repeater-laying operations in Newfoundland and at sea.

The terrain and conditions in Newfoundland were quite unlike those with which the British Post Office normally has to contend, involving trenching and cabling through bog, rock and ponds in country of which no detailed survey or maps were available. Maps were constructed from aerial survey, and alternative routes were explored on foot before a final choice was made. As much use as possible was made of water sections in the sea, river estuary and ponds; some 22 miles were accounted for in this way, leaving about 41 miles to be trenched by machine or blasted. A contractor was engaged for this purpose and to lay the cable in the trench, but all jointing was done by the Post Office. The standards of conductor and core jointing were the same as those in the cable factories

and on ship, portable injection-moulding machines and X-ray equipment being specially designed for handling over the bog. A single pair cable was also laid in the main cable trench to provide speaker facilities between Clarenville and Terrenceville (which has no public telephone), with intermediate positions for use of the lineman. As a measure of protection against lightning strikes, two bare copper wires were buried about 12 inches apart and 6 inches above the cable. Both the constructional work in Newfoundland⁴ and the laying operation at sea⁵ have been described elsewhere.

TEST RESULTS

In the interval between the completion of the link in May, 1956, and its incorporation in the transatlantic system, tests were carried out to establish its performance and day-to-day variations; an assessment of the annual variations has, of course, been impossible at this date.

Variation of Transmission Loss

Close observation of the transmission loss of the 92 kc pilots on Groups 1 and 5 leads to the following tentative conclusions:

(a) Over periods of 1 hour the variations are not measurable, i.e., less than ± 0.05 db.

(b) Over periods of 24 hours there are no systematic changes; apparently random changes of about 0.1 db are probably attributable to the measuring equipment.

(c) Over a period of eight weeks (July and August, 1956) there was a systematic increase in loss of about 0.3 db. By means of the loop-gain equipment it has been possible to deduce that most of this change has occurred in the land section.

The results indicate that the submarine cable link has better day-to-day stability than the best testing equipment which it has been possible to provide. Many more data will clearly be necessary before the annual variations can be definitely established, but the present indications are that these will be less than those assumed in the design of the link.

Attenuation/Frequency Characteristics

The frequency characteristics of the supergroup in the two directions of transmission are shown in Fig. 14. It will be seen that in no transatlantic groups does the deviation from mean exceed ± 0.35 db.

Circuit Noise

Table II shows the noise level on Channels 1 and 12 of each of the five groups measured without traffic on the system.

To assess the magnitude of intermodulation noise, all channels in one direction were loaded simultaneously with white noise and measurements taken on each channel in the opposite direction. From the talker volume data assumed in the design of the system, the expected mean talker power is -11.1 dbm at a point of zero relative level, with an activity of 25 per cent. For an equivalent system loading, therefore, the level of white noise applied to each channel under the above test conditions should be -14.1 dbm. Since this loading gave no sensible increase in the circuit noise, the test levels were raised until a reasonable increase in the noise level was obtained. In order to raise the channel noise to the specified maximum of 28 dba it was necessary to raise the channel levels to about -1 dbm and -4 dbm in the lower and upper bands respectively. These levels, some 13 db and 10 db above the assumed maximum loading of the system, give noise levels at least 26 db and 20 db above normal, and it is seen that adequate margins exist for variations and deterioration of the link.

Closely allied to the problem of intermodulation is the overload characteristic of the system. Table III shows the measured overload point of the link expressed as an equivalent level at the output of the amplifier in the repeater nearest to the transmitting terminal. It also shows the margin between the channel level at that point and the overload point of the system; according to Holbrook and Dixon¹³ the minimum requirement in this respect is 18 db.

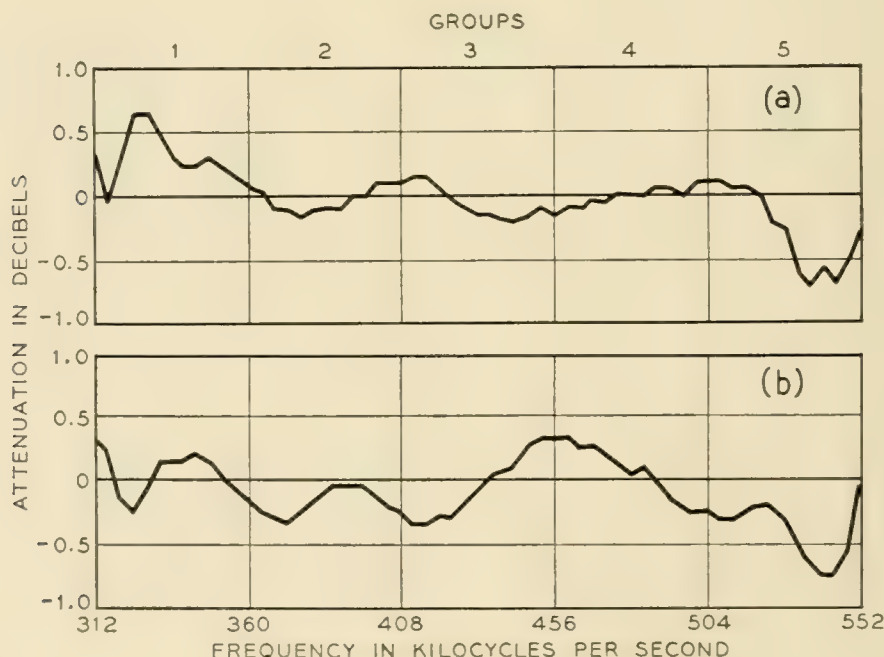


Fig. 14 — Attenuation versus frequency characteristics of supergroup. (a) Sydney Mines—Clarendville. (b) Clarendville—Sydney Mines.

TABLE II

Group	Channel	Noise level	
		Sydney Mines	Clarenville
		dba	dba
1	1	25.0	24.5
1	12	24.5	23.2
2	1	24.0	22.5
2	12	23.5	22.0
3	1	24.0	20.5
3	12	25.0	17.5
4	1	24.5	17.5
4	12	24.5	16.5
5	1	24.5	17.5
5	12	27.0	18.0

These results justify the assumption made in the design of the link, that intermodulation noise is negligible.

Crosstalk

The crosstalk requirements are met in all respects.

CONCLUSIONS

The submarine-cable link between Clarenville, Newfoundland, and Sydney Mines, Nova Scotia, was completed in May, 1956, and provides five carrier telephone groups, each capable of carrying twelve high-grade telephone circuits or their equivalent. The transmission objectives have been met in every respect.

Three 12-circuit groups are connected to the three groups across the Atlantic between Scotland and Newfoundland; the other two groups are available to provide 24 circuits between Newfoundland and the mainland of Canada.

TABLE III

Frequency	Equivalent at amplifier in first repeater		
	Channel level	Overload	Margin
kc	db	db	db
552	-2	+20	22
312	-4	+24	28
260	-5	+25	30
20	-5	+25	30

ACKNOWLEDGMENTS

It has been the authors' privilege to present an integrated account of the work of many of their colleagues in the Post Office and in industry. Post Office staff have been responsible for designs and for inspection and testing at home and in the field, as well as the laying of the submarine-cable system by H.M.T.S. *Monarch*. In Great Britain, Submarine Cables, Ltd., and the Southern United Telephone Co., Ltd., provided the submarine and overland cables respectively, while Standard Telephones and Cables, Ltd., supplied and contributed much to the design of the submerged repeaters and terminal equipment. On site, the assistance rendered by the Ordnance Survey of Great Britain, the Canadian Comstock Co., Ltd., who laid the cable across Newfoundland, the Northern Electric Co., Ltd., who carried out the terminal equipment installations, and by the other partners in the project, the American Telephone and Telegraph Co. Inc., the Canadian Overseas Telecommunication Corporation and the Eastern Telephone and Telegraph Co., has been invaluable. The permission of the Engineer-in-Chief of the Post Office to make use of the information contained in the paper is gratefully acknowledged.

11. BIBLIOGRAPHY

1. Transatlantic Cable Construction and Maintenance Contract, Nov. 27, 1953.
2. E. T. Mottram, R. J. Halsey, J. W. Emling and R. G. Griffith, Transatlantic Telephone Cable System — Planning and Over-All Performance. See page 7 of this issue.
3. M. J. Kelly, Sir Gordon Radley, G. W. Gilman and R. J. Halsey, A Transatlantic Telephone Cable, Proc. I.E.E., **102B**, p. 117, Sept., 1954, and Communication and Electronics, **17**, pp. 124-136, March, 1955.
4. H. E. Robinson and B. Ash, Transatlantic Telephone Cable — The Overland Cable in Newfoundland, Post Office Electrical Engineers' Journal, **49**, pp. 1 and 110, 1956.
5. J. S. Jack, Capt. W. H. Leech and H. A. Lewis, Route Selection and Cable Laying for the Transatlantic Cable System. See page 293 of this issue.
6. H. A. Affel, W. S. Gorton and R. W. Chesnut, A New Key West-Havana Carrier Telephone Cable, B.S.T.J., **11**, p. 197, 1932.
7. R. J. Halsey and F. C. Wright, Submerged Telephone Repeaters for Shallow Water, Proc. I. E. E., **101**, Part I, p. 167, Feb., 1954.
8. R. J. Halsey, Modern Submarine Cable Telephony and Use of Submerged Repeaters, Jl. I. E. E., **91**, Part III, p. 218, 1944.
9. A. W. Lebert, H. B. Fischer and M. C. Biskeborn, Cable Design and Manufacture for the Transatlantic Submarine Cable System. See page 189 of this issue.
10. R. A. Brockbank, D. C. Walker and V. G. Welsby, Repeater Design for the Newfoundland-Nova Scotia Link. See page 245 of this issue.
11. G. H. Metson, E. F. Rickard and F. M. Hewlett, Some Experiments on the Breakdown of Heater-Cathode Insulation in Oxide-Coated Receiving Valves, Proc. I. E. E., **102B**, p. 678, Sept., 1955.
12. J. F. P. Thomas and R. Kelly, Power-Feed System for the Newfoundland-Nova Scotia Link. See page 277 of this issue.
13. B. D. Holbrook and J. T. Dixon, Load Rating Theory for Multi-Channel Amplifiers, B.S.T.J., **18**, p. 624, 1939.

Repeater Design for the Newfoundland-Nova Scotia Link

By R. A. BROCKBANK,* D. C. WALKER* and
V. G. WELSBY*

(Manuscript received September 15, 1956)

The Newfoundland-Nova Scotia cable required the provision of 16 submerged repeaters each transmitting 60 circuits in the bands 20–260 kc from Newfoundland to Nova Scotia and 312–552 kc in the opposite direction. The paper deals with the design and production of these repeaters. Each repeater has a gain of 60 db at 552 kc, and the amplifier consists of two forward amplifying paths with a common feedback network. Reliability is of paramount importance, and production was carried out in an air-conditioned building with meticulous attention to cleanliness and to very rigid manufacturing and testing specifications. The electrical unit is contained in a rigid pressure housing 9 feet long and 10 inches in diameter with the sea cables connected to an armor clamp and a cable gland at each end. A submerged equalizer was provided near the middle of the sea crossing.

INTRODUCTION

The British Post Office has engineered many shallow-water submerged-repeater systems,¹ and there has been a progressive improvement in design techniques and in the reliability of components which has been reflected in a growing confidence in the ability to provide long-distance systems having an economic life. The seven-repeater scheme from Scotland to Norway laid in 1954 introduced for the first time repeaters which would withstand the deepest ocean pressure together with an electrical circuit which embodied improved safety and fault-localizing devices. Also, since a repeater is only as reliable as its weakest component, much greater attention and control was directed at this stage to the design, manufacture and inspection of all components, both electrical and mechanical. This repeater design was, in fact, envisaged as a prototype for a future transatlantic project.

* British Post Office.

In the finalized plans² for the Newfoundland-Nova Scotia cable it was required to carry 60 circuits, so that some redesign of the 36-circuit 'prototype' repeater became essential. It was accepted, however, as a guiding principle throughout the redesign that there should be no departure from previous practice without serious consideration and adequate justification. It was obviously not an occasion to experiment with new ideas.

Post Office and Bell Telephone Laboratories experiences were pooled for the project, and the whole technical resources of both organizations were freely available at all times for consultative purposes. Detailed manufacturing and testing specifications were exchanged and approved, and each party was free to inspect the other's production methods. This mutual interchange was undoubtedly highly beneficial, and in the British case it resulted in a still more rigorous control of manufacturing and inspection methods.

PLANNING

General

Preliminary design calculations indicated that it should be possible to increase the circuit-carrying capacity of the Anglo-Norwegian prototype repeater from 36 to 60 circuits, and tests on a model confirmed that, with band frequencies of 20–264 and 312–552 kc, a 60.0 db gain at 552 kc could be realized with satisfactory margins against noise, distortion and overload. This gain fixed the repeater spacing at about 20.0 nautical miles so that on the selected route two repeaters would be required on the land section between Clarendville and Terrenceville and 14 in the Terrenceville-Sydney Mines sea section. It was noted that the land repeaters might have to work with an ambient temperature 12°C higher than in the sea repeaters.

Each repeater would need to be energized with a direct current of 316 ma at 124 volts so that the total route voltage would be about 2,300 volts. This voltage would be quite acceptable to the repeaters, but for normal operation it was proposed² to feed from both terminals simultaneously, thereby halving the maximum voltage to ground.³

The precise localization of any faulty or aging repeater would be of paramount importance. It was decided to retain the two supervisory methods which on the prototype had worked satisfactorily in this respect. These consisted of a pulse-distortion equipment requiring no additional components in the repeater and a loop-gain monitoring set involving a special unit in the repeater and the allocation of a 4-kc band (260–264 kc) for its operation.

Manufacture and testing of the electrical units was again to be carried out by a contractor in a temperature- and humidity-controlled production building, and in order to enable manufacture to start as early as possible, arrangements were made for the contractor to co-operate with the Post Office at an early design stage so that engineering could follow fast on the heels of the design. The outer housing and method of brazing-in the bulkheads had all proved entirely satisfactory on the prototype, and therefore these operations could proceed according to previous production. The Post Office assumed responsibility for the production and testing of the glands, since no contractor had experience of this work.

Forward planning in early 1954 scheduled the first electrical unit to be completed in June, 1955, with units following at 5-day intervals. This target was, in fact, delayed until August, 1955, but all 16 working repeaters were available fully tested before the commencement of the alying operation in May, 1956.

Distortion Monitoring Equipment

The pulsed-carried supervisory method as used on previous systems¹ is employed primarily for measuring the distortion on repeaters. Under normal operating conditions the distortion level may be only just noticeable above the noise, but should appreciable distortion occur, e.g., failure of one amplifying path in a repeater, it could be readily located, since the pulse amplitude from the faulty repeater would increase by about 12 db for second-harmonic distortion and about 18 db for third-order distortion.

Loop-Gain Monitoring Equipment

For the loop-gain monitoring equipment¹ the repeaters have to be designed to incorporate a second-harmonic generator operating at a frequency unique to each repeater at 120-cycle spacing in the 260–264-kc band. The second harmonics return to Sydney Mines in the 520–528-kc band, and these two channels must be taken out of service during the measurement.

Levels and Equalization

Controlling Factors.

In practice, deviations from an ideal system wherein all repeaters match the cable and operate at all times at the same levels require the repeaters to be designed with specific margins against overload and intermodulation to meet an agreed maximum noise figure for the system

under all working conditions.² Factors involved in assessing these margins and in planning the equalization and level diagram for the system are as follows:

(a) *Temperature*. — The final assumed sea-bottom temperature was 2.3°C , with a maximum annual variation of $\pm 3^{\circ}\text{C}$. The maximum change in attenuation might therefore be ± 4 db at 552 kc. The land section change would be ± 3 db at 552 kc due to a possible $\pm 10^{\circ}\text{C}$ change on a mean of 7.5°C . The effect of these seasonal changes would be reduced by the provision of manually adjusted equalization at Clarendville, Terrenceville and Sydney Mines.

The repeaters show a small change in gain (less than 0.05 db) during the warming-up period after energization, but the effect of ambient-temperature change is negligible.

(b) *Repeater spacing*. — The repeater-section cable lengths were to be cut in the cable factory such that the expected attenuation at 552 kc when laid at the presumed mean annual temperature of the location should be 60.0 db. An anticipated decrease in attenuation of 1.42 per cent at 552 kc was assumed when laid. The assumed mean annual temperature of sections of the route varied between 1.7 and 4.0°C . Temperature corrections employed an attenuation coefficient at 552 kc of $+0.16$ per cent per degree centigrade. It was expected that the total error at 552 kc after laying seven repeaters would not exceed 1.5 db, and this could be largely corrected as explained in (c).

(c) *Cable Characteristics*. — The cable equalization built into the repeater was based on a cable attenuation characteristic which was later discovered to be appreciably different from the laid characteristic. Cutting the cable as described in (b) overcomes this difficulty at 552 kc, where the signal/noise ratio is at a minimum. The new shape of the characteristic, however, indicated that at about 100 kc the error would reach 7 db on the complete route. To reduce this deviation it was decided to introduce a submerged equalizer in the middle of the sea section to correct for half this error and to insert in each of the four-wire paths of the transmit and receive equipments equalization for one-quarter of this error. There is an appreciable signal/noise margin in hand at this frequency, so that the system would not be degraded below noise specification by these equalizer networks.

It was also decided that the splice at the equalizer which would connect the halves of the link together should not be completed until after the laying operation had commenced. An excess length of cable was provided on the equalizer tail, and this could be cut at a position indicated by measurements taken during the laying of the first half-section so that

the equalization at the 552-kc point could be largely corrected for laying and temperature-coefficient errors. It is not, in practice, easy to separate these two factors.

(d) *Repeater characteristic.* — The repeater was designed to equalize the original cable-attenuation characteristic to ± 0.2 db, as this was possible with a reasonable number of components. This variation appeared as a roll in the gain/frequency characteristic, which was expected to be systematic and would therefore lead to a ± 3 db roll in the overall response. It was proposed that equalization for this should be provided at the receive terminal. Manufacturing tolerances were expected to be small and random.

(e) *Repeater interaction.* — At the lower frequencies where the loss of a repeater section is comparatively small, a roll in the overall frequency response will arise due to changes in the interaction loss between repeaters. The design aimed at providing a loop loss greater than 50 db which would reduce rolls to less than ± 0.03 db per repeater section and therefore to about 0.5 db at 20 kc with systematic addition on the whole route.

Planning of Levels

From a critical examination of all these variables it was concluded that the repeater should be designed to have an overload margin of 4 db above the nominal mean annual temperature condition. It was also desirable for the system to be able to operate within its noise allowance if one path of a twin amplifier failed. Tests on a model amplifier gave overload values of +24 dbm and +19 dbm for two- and one-path operation, respectively, so that with a single-tone overload requirement of 18 dbm⁴ at a zero-level point, the maximum channel level at the amplifier output would be -3 dbr for a single amplifying path.

Thermal-noise considerations (i.e. resistance plus tube noise) fixed the minimum channel level at the repeater input at -69 dbr in order to meet the allowable system noise limit of +28 dba at a zero-level point. At 552 kc the amplifier gain is 65 db, so that the minimum level at the amplifier output is -4 dbr. A system slope of ± 4 db due to temperature variations, corrected by similar networks at the transmit and receive terminal, would, however, degrade the noise by 0.5 db. Intermodulation noise was estimated⁵ on an average busy-hour basis, and it was concluded that the increase in noise at 552 kc from this source was negligible — less than 1 db, even with several repeaters in which the amplifier had failed on one path. At lower frequencies the contribution from intermodulation noise is greater, and at 20 kc it exceeds resistance noise.

However, at 20 kc the total noise is some 8 db below the specification limit, and therefore again several amplifiers could fail on one path before the noise exceeded the specification limit. Actually it was discovered that the predominant source of third-order intermodulation on the repeater was in the nickel-iron/ceramic seals on high-voltage capacitors and followed a square law with input levels.

From a more detailed examination of the factors briefly mentioned above it was decided that the initial line-up should be based on a nominal flat -3.5 dbr point at the amplifier output and the final working levels decided upon as the results of tests on the completed link.

With equal loading on the grid of the output tube at all frequencies the worst signal/noise ratio exists at 552 kc; some pre-emphasis of the transmit signal should therefore prove to be beneficial. In fact, after completing the tests on the link it was decided to improve the margin on noise by raising the level at 552 kc by 2 db, thus giving a sloping level response at the amplifier output in the high-frequency band. To maintain the same total power loading, the low-frequency band levels were decreased by 1 db, still retaining a flat response.

Laying

It was proposed to use laying methods with continuous testing similar to those employed successfully on the Anglo-Norwegian project. The complete link with a temporary splice at the equalizer would be assembled and tested on board H.M.T.S. *Monarch* and laying would proceed from Terrenceville to Sydney Mines in the high-frequency direction of transmission. A detailed description of the actual laying operation is given elsewhere.⁶ After completion of tests on the submarine section the land section to Clarendville would be connected with appropriate equalization at Terrenceville.

DESIGN OF ELECTRICAL UNIT OF SUBMERGED REPEATER

General

The equipment is contained in a hermetically sealed brass cylinder (filled with dry nitrogen) $7\frac{3}{4}$ inches in diameter and 50 inches long, which is bolted at one end to one of the bulkheads of the housing. A flexible coaxial cable emerges through an O-ring seal at each end, and these are ultimately jointed to the cable glands. The various units forming the complete electrical unit are mounted within a framework of Perspex (polymethylmethacrylate) bars which forms the main insulation of the

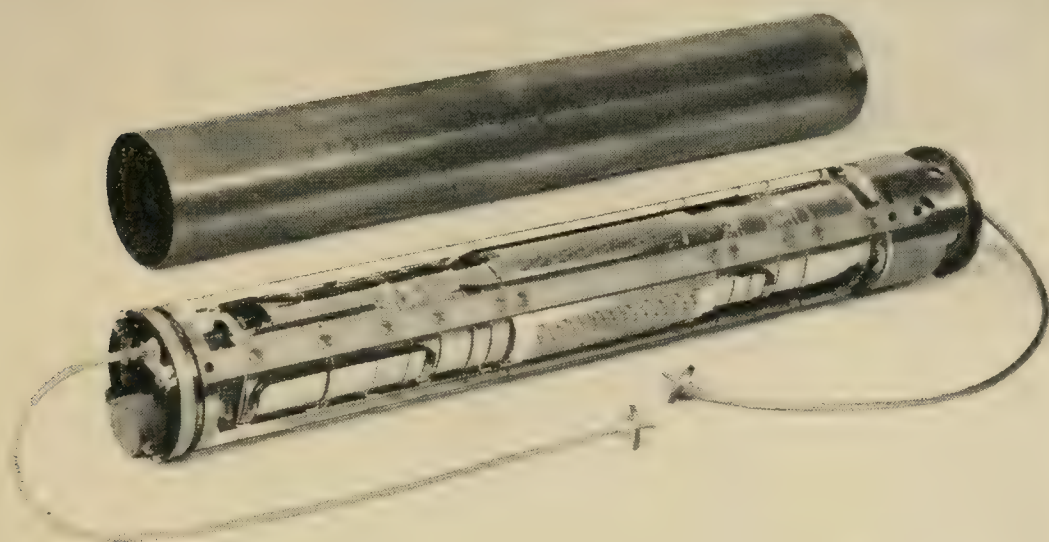


Fig. 1 — Internal unit.

repeater, and these units may operate at 3 kv dc to the grounded brass cylinder. Fig. 1 shows the construction.

A schematic of the electrical circuit is given in Fig. 2. The direct current for energizing the repeater is separated from the carrier transmission signals by the A- and B-end power separating filters, and passes through the amplifier tube heaters and a chain of resistors developing 90-volt high-voltage supply for the amplifier. The carrier-frequency signals pass through the same amplifier via directional filters. Equalization is provided in the amplifier feedback circuit (about 20 db) and in the equalizers and the bridge networks which combine the directional filters. The main purpose of the bridges, however, is to reduce the severe harmonic requirement on the directional filters due to having high- and low-level signals present at the repeater terminals. The whole carrier circuit is designed on a nominal impedance of 55 ohms. Attached to the B-end of the repeater is the loop-gain supervisory unit and also, via a high-voltage fuse, a moisture-detector unit used primarily during the high-pressure test to confirm that the housing is free from leaks. The latter comprises a series-resonant circuit at about 1.3 mc, in which the inductance is varied by the gas pressure on an aneroid capsule mounted in the space between the electrical unit and the housing. The presence of moisture in this cavity increases the gas pressure owing to the release of hydrogen by the reaction of water vapor with metallic calcium held in a special container. At a later stage the fuse is blown to disconnect this circuit.

The circuit design of the repeater introduces multiple shunt paths

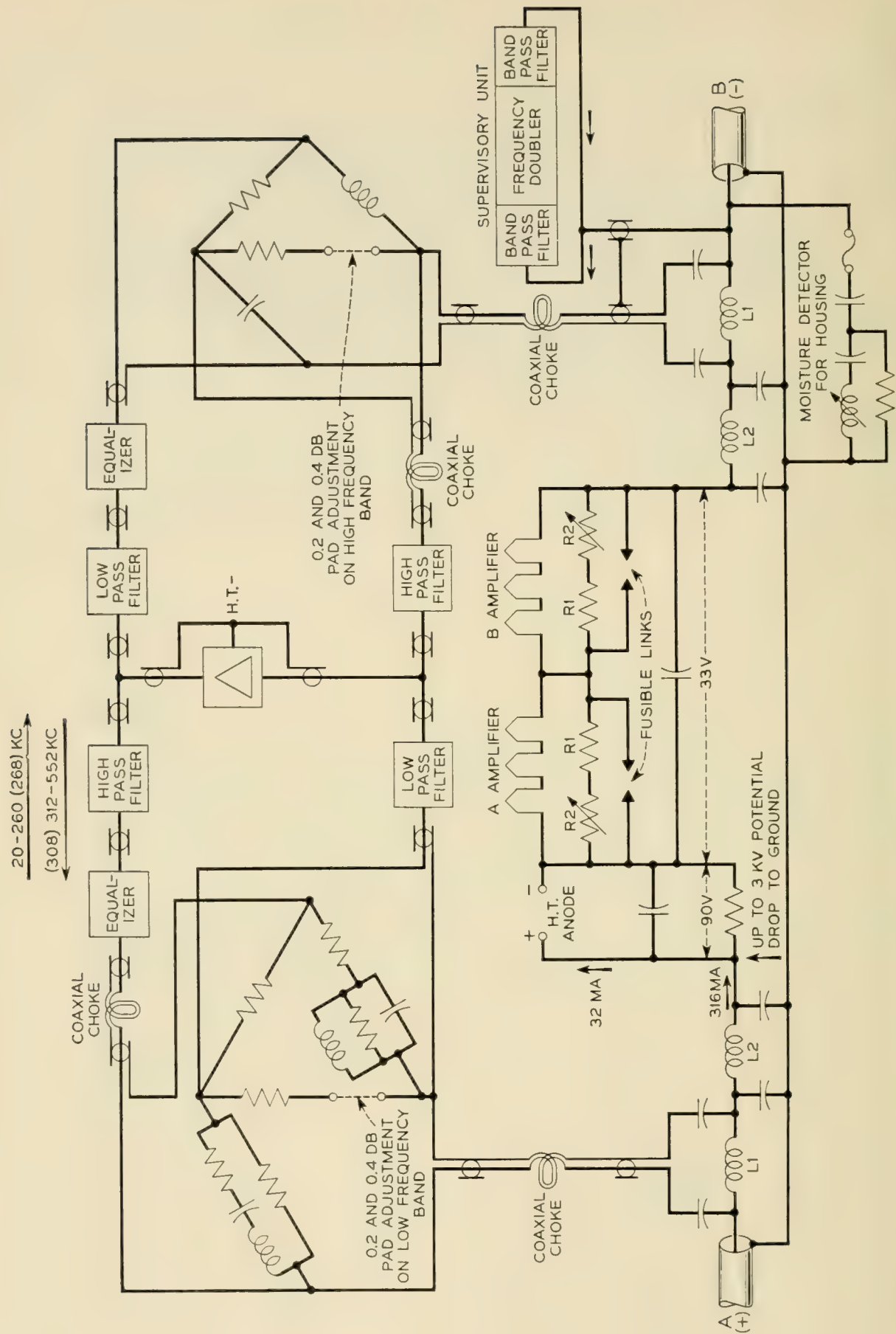


Fig. 2 — Schematic of submerged repeater.

across the amplifier, and care has to be taken to ensure that there is adequate attenuation in each path. In general, the design is such that the combination will give a loop loss of at least 40 db in the working band (to reduce rolls in the gain characteristic) and 20 db at all frequencies (as a guard against instability), even when one repeater terminal is open- or short-circuited to simulate a faulty cable.

Unit Details

Power Filter.

The power filters are, in effect, a series pair of high- and low-pass filters (see Fig. 2). The shunt capacitors may have to withstand 3 kv, and clearances on the input cable and some wiring have to be adequate for this voltage. The inductors have to carry the line current of 316 ma dc, and the intermodulation must be extremely low (see section on inductors below).

Directional Filter.

The directional filters are a conventional Zobel high-pass and low-pass filter pair with a susceptance-annulling network. Silvered-mica capacitors and carbonyl-iron dust-cored inductors are used. The bridges combining the 'go' and 'return' filters reduce the distortion due to the ferromagnetic material to an acceptable level.

Bridge and Equalizer.

A simple non-resonant bridge is used at the B-end of the repeater, but the A-end bridge is a resonant type and provides a substantial degree of equalization (see Fig. 2).

The equalizers are of conventional form. Trimming capacitors (selected on test) were provided for critical capacitances in order to utilize standard tolerances on all capacitors. A pad of 0.2 db and 0.4 db is provided on each equalizer unit so that the repeater low-frequency or high-frequency path can be independently trimmed to give the best match to the target response for the repeater.

The components in the above circuit were small air-cored inductors, silvered-mica capacitors, and wire-wound resistors, except for a few high-resistance ones, which were of the carbon-rod type. Included in this unit are coaxial chokes whose purpose is to separate parts of the circuit to avoid the effect of multiple grounding. They are merely inductors wound with coaxial wire on 2-mil permalloy C tape ring cores.

Supervisory Unit.

The supervisory unit comprises a frequency-selection crystal filter of about 100-cycle bandwidth in the range 260–264 kc fed from the low-frequency output end of the repeater via a series resistor. This filter feeds a full-wave germanium point-contact crystal-rectifier bridge which acts as a frequency doubler. The second harmonic in the band 520–528 kc is filtered out by a coil-capacitor band-pass filter, and fed back through a resistor to the same point in the repeater. The two series resistors minimize the bridging loss of the unit on the repeater and ensure that a faulty supervisory component has negligible effect on the normal working of the repeater.

DC Path.

The dc path includes a resistor providing the 90-volt supply and the heater chain of six electron tubes (see Fig. 2). The voltage drop across the heater chain is not utilized for the amplifier high-voltage supply, as the heaters would then be at a positive potential with respect to the cathodes, thereby increasing the risk of breakdown of heater-cathode insulation. There would also be a complication in maintaining the constant heater current, particularly should the high-voltage supply current fail in one path of the amplifier. The normal amplifier high-voltage supply current is 32 ma.

It is essential to maintain a dc path through the repeater even under fault conditions in order that fault-location methods can be applied. Special care has therefore been taken to provide parallel paths capable of withstanding the full line current. For example, the high-voltage resistor actually consists of a parallel-series combination of ten resistors, and the whole assembly is supported on Sintox (a sintered alumina) blocks which maintain a good insulation at 3 kv dc, even at high temperatures.

Electron tube operation for consistent long life indicates the necessity to maintain a specific constant cathode temperature, and to achieve this, electron tubes are grouped according to heater characteristics into six heater-current groups between 259 and 274 ma and stabilized to ± 1 per cent. The appropriate heater-shunt resistor is applied so that the tube operates correctly with 316-ma line current, but for convenience the shunt is taken across each set of three tubes, all in one heater group, forming one amplifier path. R1 is fixed (300 ohms) and R2 is selected to suit the tubes. R1 is the resistance winding of a special short-circuiting fuse; when energized by the full line current should a heater become

open-circuited, it causes a permanent direct short-circuit across the heater chain. The line voltage will be temporarily increased by about 95 volts while the fuse operates (1 min) and will then drop to 12 volts below normal.

Amplifier.

The amplifier circuit is shown in Fig. 3. It consists of two 3-stage amplifiers connected in parallel between common input and output transformers with a single feedback network. This circuit arrangement allows one amplifier path to fail without appreciably affecting the gain of the complete amplifier (less than 0.1 db for all faults except those on the grid of V1 and the anode of V3, but the overload point is reduced by about 5 db and distortion at a given output level is increased (about 12 db for second harmonics). Care has been taken to ensure that the open- or short-circuiting of a component in one amplifier path will not affect the performance, life or stability of the remaining path, and this involves the duplication of certain components.

Mixed feedback is employed to produce the required output impedance; the current feedback is obtained from the resistor feeding the high-voltage supply to the output transformer, and the voltage feedback is developed across a two-turn winding on the output transformer, which also serves as a screen. The output of the feedback network is fed in series with the input signal to the grid of V1. The gain response of the amplifier is chiefly controlled by the series-arm components in the feedback network, which resonate at 600 kc.

The input transformer is built out as a filter and steps up in impedance from 55 ohms to 17,000 ohms. Protective impedances minimize the effect of a short-circuit on the grid of one of the first-stage tubes. The anode load of the first stage resonates at 600 kc, and is roughly the inverse of the feedback network so as to give constant feedback loop gain over the working frequency band. The output tube has about 5.5 db of feedback from its cathode resistor, and the pair of output tubes feed the output transformer, which steps down from 5,000 to 55 ohms.

Specially designed long-life tubes are used.⁷ The first two stages are operated at about 40 volts on the screens and anodes; each anode current is 3 ma, giving a mutual conductance of 5.1 ma/v. The output stage is operated at 60 volts on the screen, +15 volts on the suppressor grid to sharpen the knee of the V_a/I_a characteristic, and nearly the full high-voltage supply of 90 volts on the anode; the anode current is 6 ma, giving a mutual conductance of 6.6 ma/v. The tube dynamic impedance is approximately 300,000 ohms. To obtain an anode current nearest to the

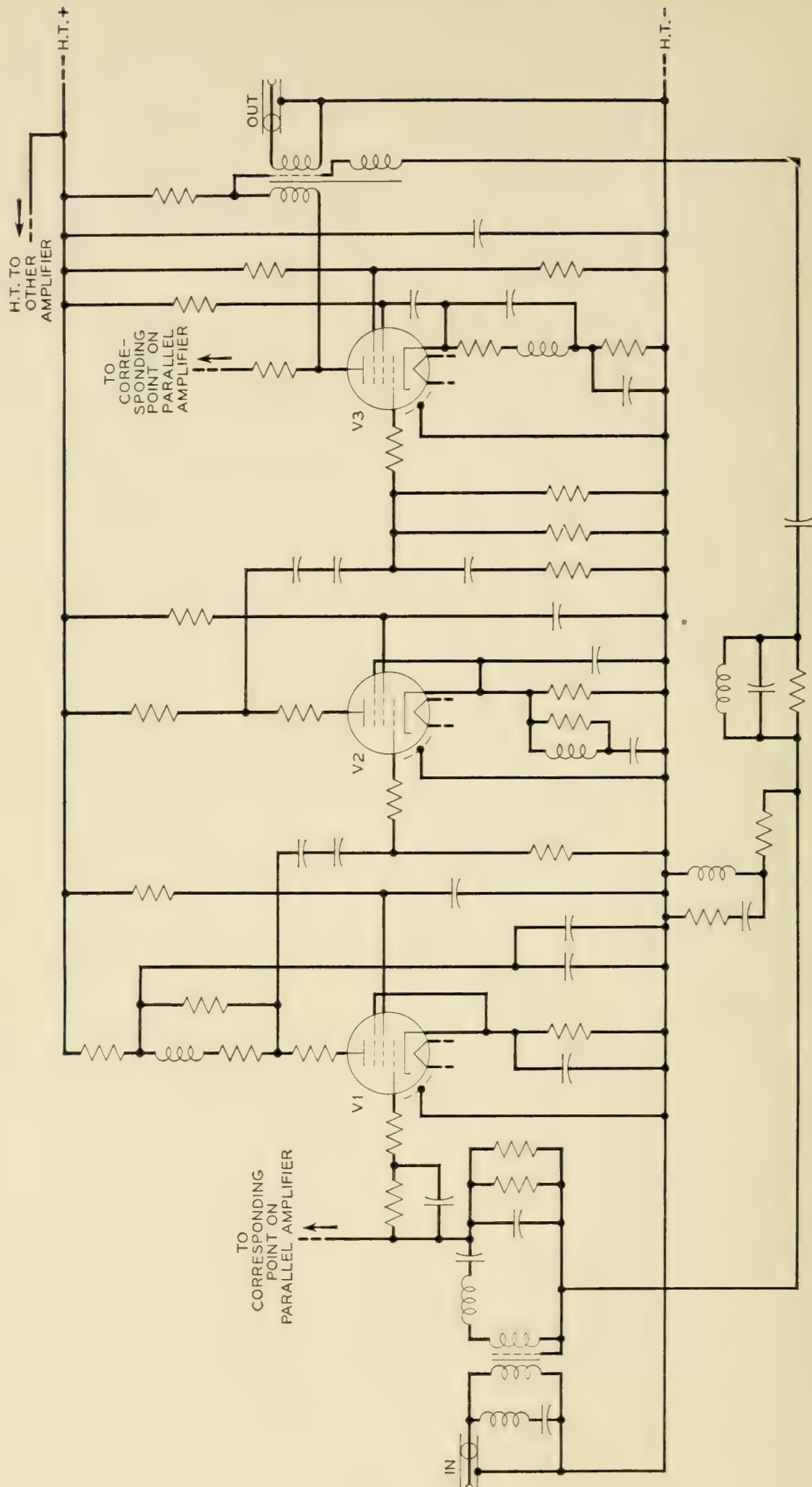


Fig. 3 — Amplifier circuit.

design value (and for which the tubes are aged), one of two values of bias resistance can be selected for V1 and V2, and one of three values of bias resistance for V3.

All capacitors subject to the high-voltage supply voltage are of the oil-filled paper type, and the others are of the silvered-mica type. Inductors are air-cored spools which are multi-sectioned when used in high-impedance circuits. All resistors are of the solid carbon-rod type, except for the input-transformer termination, which consists of two high-stability cracked-carbon resistors, and those in the feedback network, which are wire wound.

The input and output transformers employ 2-mil permalloy C laminations, and the latter core is gapped on account of the polarizing current. A narrow Perspex spool fits the center limb, and conventional layer windings are used; the screen is a sandwich made of copper foil with adhesive polythene tape.

3.3 *Mechanical Design Details*

The arrangement can be seen from Fig. 1. At the A-end is a cast-brass pot containing the resistors providing the amplifier high-voltage supply, and as this is bolted directly on to the housing bulkhead, the heat generated is readily conducted away. The remainder of the units are in cylindrical cans mounted in the insulating framework formed by four Perspex bars. These are sprayed with copper on both faces to guarantee the dc potential on these surfaces and eliminate the risk of ionization at working voltages. The cans are not hermetically sealed but are dried out with the repeater when it is finally sealed and filled with dry nitrogen. Perforated covers on the amplifier allow air circulation to reduce the ambient temperature.

Fig. 4 shows a typical can assembly, and Figs. 5 and 6 the construction of the amplifier. It will be seen that the latter is a double-shelved structure with tubes alternating in direction, and an amplifier path is located on each side of the chassis; the input and output transformers are at opposite ends, and the feedback network is contained in a hermetically sealed can in the center of the unit.

All cans are finished with a gold flash which is inert and gives a clean appearance stimulating a high standard of workmanship. Tin plating was formerly used, but it has been shown that tin tends to grow metallic whiskers.⁸ Unfortunately certain capacitor cans had to be tin plated, and extra precautions consisting of wide clearances or protective shields have had to be taken. The risks from growth on soldered surfaces is not

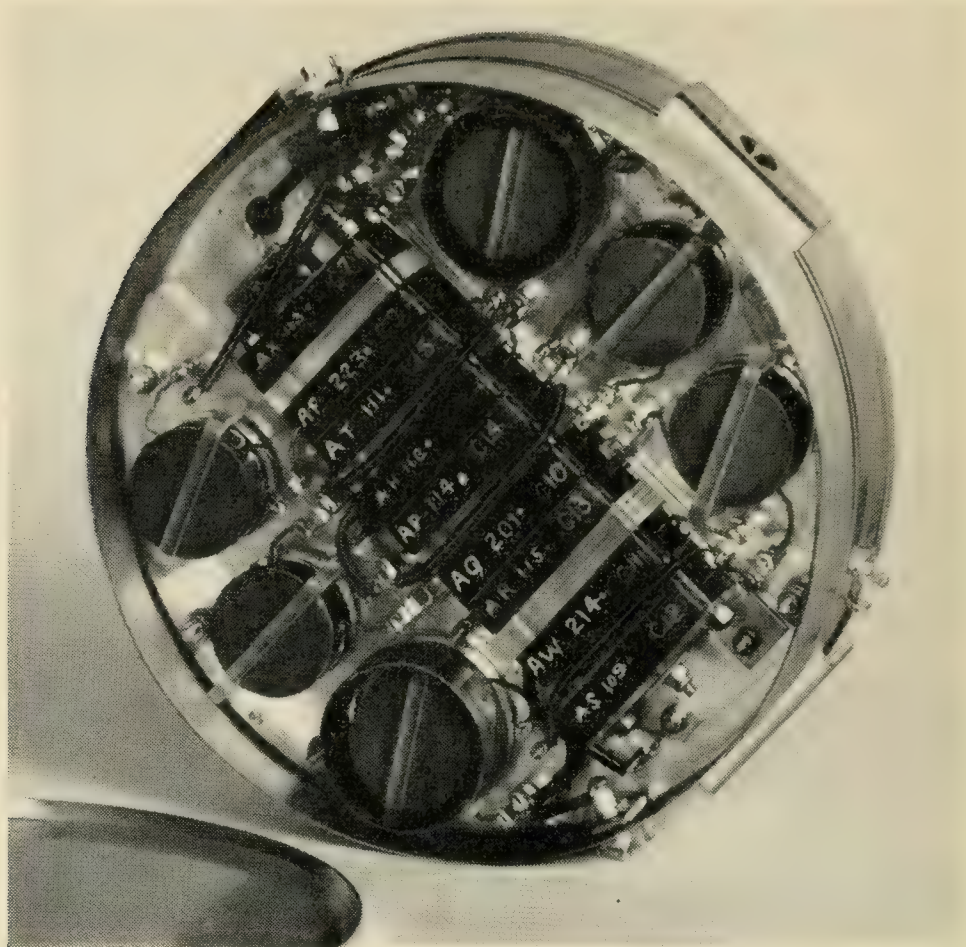


Fig. 4 — Directional filter unit. High-pass filter with low-pass filter can on rear.

thought to be great, as all solders used have less tin content than the eutectic alloy. All connecting wires are gold plated instead of the usual tin plating.

The insulating sub-panels in units are usually made of Perspex, but where the items are subjected to high temperatures (e.g., resistance box), Sintox, a sintered alumina ceramic, or Micalox is used.

Polytetrafluoroethylene (p.t.f.e.) is another insulant used, and p.t.f.e.-covered wire threaded through copper tube forms the coaxial interconnecting leads between the can units.

Careful attention is paid to the mounting of components. Small resistors, etc., are supported by soldering to tags which are the appropriate distances apart, and multi-limb tags are employed to minimize the number of soldered joints. Where it is essential to solder more than one wire per limb on a tag, they must be soldered at the same time.

Larger components are clamped. Electron tubes are mounted in a holder so as to facilitate preliminary testing with 'standard electron tubes,' and they have a sprung nylon retainer; the final electrical connection is made by soldering on to an extension of the wire leading through the pins.

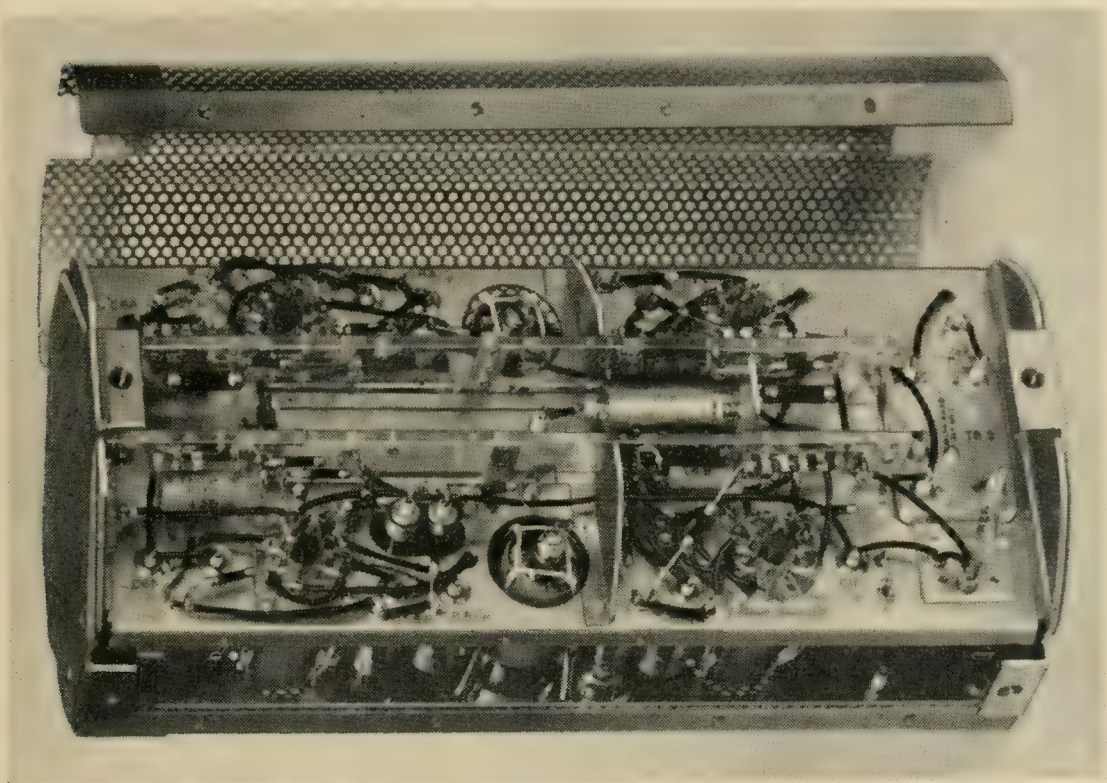


Fig. 5 — Amplifier.

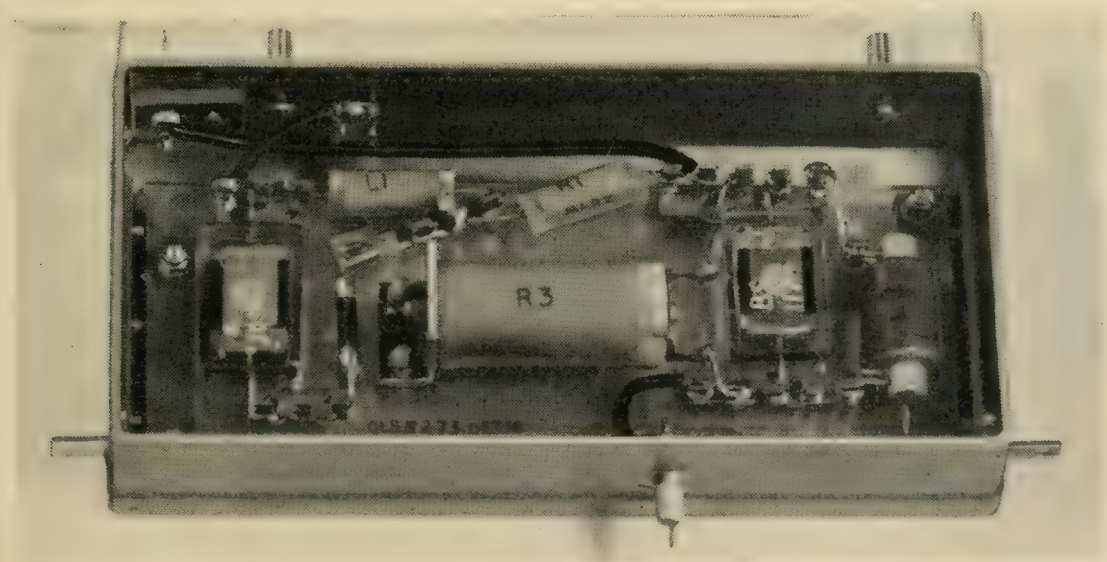


Fig. 6 — Feedback unit of amplifier.

DESIGN OF INTERNAL UNIT OF SUBMERGED EQUALIZER

The submerged equalizer corrects for the difference between the assumed design cable characteristic and the subsequently determined laid characteristic for equal attenuation lengths at 552 kc. It also absorbs the loss of 9 nautical miles of cable and has an attenuation of 26.0 db at 552 kc.

The construction is identical with the submerged repeater except for the replacement of all can assemblies, other than the power filters, by the equalizer cans. Fig. 7 is a schematic of the unit.

ELECTRICAL COMPONENTS

General

The components used in the repeaters were either designed specially for submerged repeaters or were standard items with improvements. There are approximately 300 components in each repeater of which 110 are in the amplifier. Rigorous control of manufacture and meticulous inspection is imperative to ensure a consistent long-life product, and cleanliness is essential at all stages. In some cases 'belt and braces' tech-

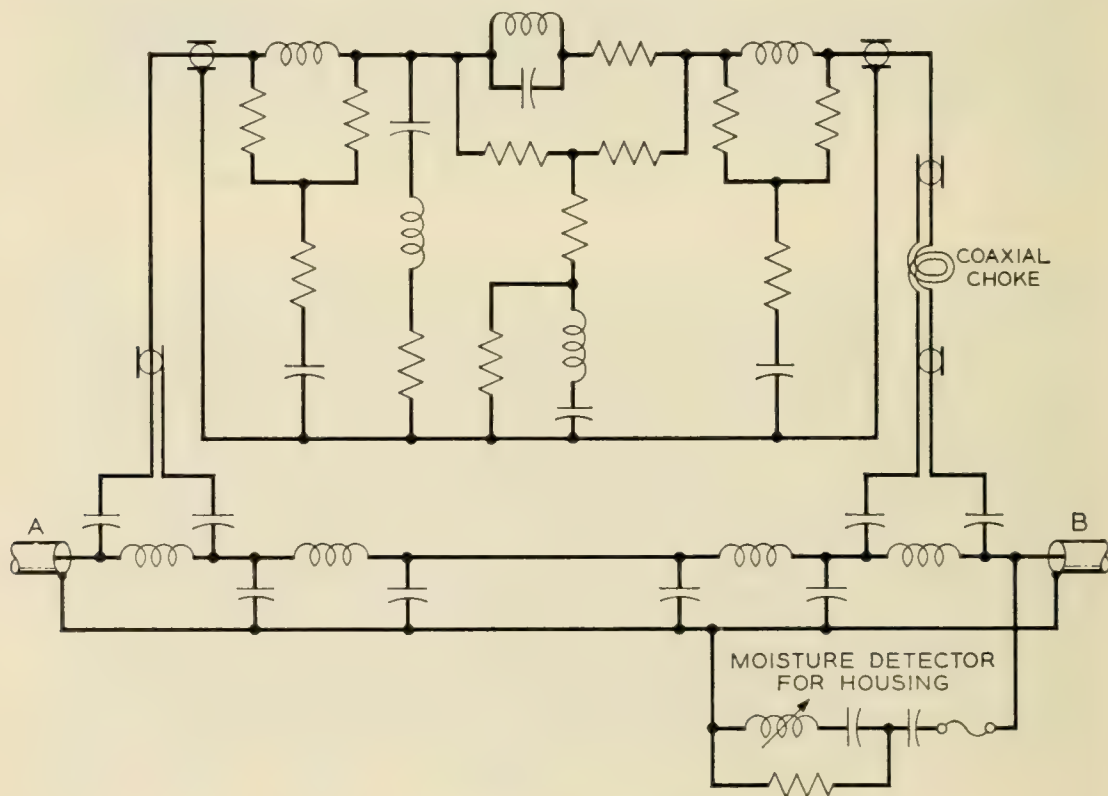


Fig. 7 — Schematic of submerged equalizer.

niques can be effectively employed, e.g., by using double connections. Over 1,500 separate soft-soldered connections are involved in the complete assembly. Much work has been done on components, but only a brief indication can be included here. A range of typical components appears in Figs. 4–6.

Resistors

Resistors fall into the following categories:

(a) Power resistors used solely for dc purposes (e.g., resistors providing the amplifier high-voltage supply). These are wire-wound vitreous-enamelled resistors on Sintox ceramic formers. Nichrome terminal leads are used, and all connections are brazed.

(b) High-frequency resistors whose tolerance is not close, and often carrying direct current but of low power (e.g., anode load resistances). A modification of a standard carbon-rod resistor is used. The ends of the rod are copper plated and the end caps and terminal leads are soldered on. The tolerance is normally ± 5 per cent, and the maximum rating permitted is about one-quarter of the commercial rating.

(c) Precise high-frequency resistors of resistance below 1,000 ohms (e.g., feedback components). Here wire-wound spool resistors are suitable, and bifilar or reverse layer windings with Lewmex enamel and silk-covered wire are used.

(d) Precise high-frequency resistors of high resistance. For terminating the input transformer a resistance of 17,000 ohms is required. Because it is not possible to make a suitable wire-wound resistor, high-stability cracked-carbon film resistors are used, but to minimize the effect of a disconnection two are used in parallel.

Inductors

The majority of inductors used in the amplifier and equalizer do not require a high Q-factor. They are wound on air-cored ceramic bobbins of four types, and the high-inductance ones are sectionalized. In general solid wire with Lewmex enamel and double-silk covering is used for the amplifier inductors, and stranded wire for the equalizer inductors.

A high Q-factor inductor is essential in the directional filters, and a carbonyl-iron pot core was used; the Q-factor is about 250 at 300 kg. Precise adjustment and stability of inductance was obtained by setting the gap between the halves of the pot core with a cement of Araldite (an epoxy resin) and titanium dioxide. A Perspex former was used.

Special inductors were required in the power filters to take the 316-ma

dc line current. A wave-wound air-cored coil was used for the carrier-path filter (L_1 in Fig. 2), and a toroid on an a.f. Permalloy dust-core ring for the low-pass filter (L_2 in Fig. 2). Solid wire, with Lewmex and double-silk covering, was used to keep the dc resistance to a minimum.

5.4 Capacitors

Capacitors are divided into three categories:

(a) Those subjected to the full line voltage which may operate at up to 3 kv.

(b) Those subjected to the amplifier high-voltage supply of 90 volts.

(c) Those which have negligible polarization (less than 10 volts), and which are often required to precise values.

Groups (a) and (b) are of the oil-filled paper type, with, respectively, four layers of 36-micron and three layers of 7-micron Kraft paper. The oil is a mineral type loaded with 18 per cent resin, and the capacitors are filled at 60°C and sealed at room temperature.

Small capacitors and those of precise value as in group (c) are silvered-mica capacitors. These are encased in an epoxy resin to give mechanical protection and a seal against moisture. Visual inspection of all mica plates is made before and after silvering, and any with cracks, inclusions, stains or any other abnormality are rejected. Mica is a very variable material and at times the percentage rejects were high, but probably many of the reasons for rejection would not have been significant as far as the life of the capacitor is concerned. However, experience has shown

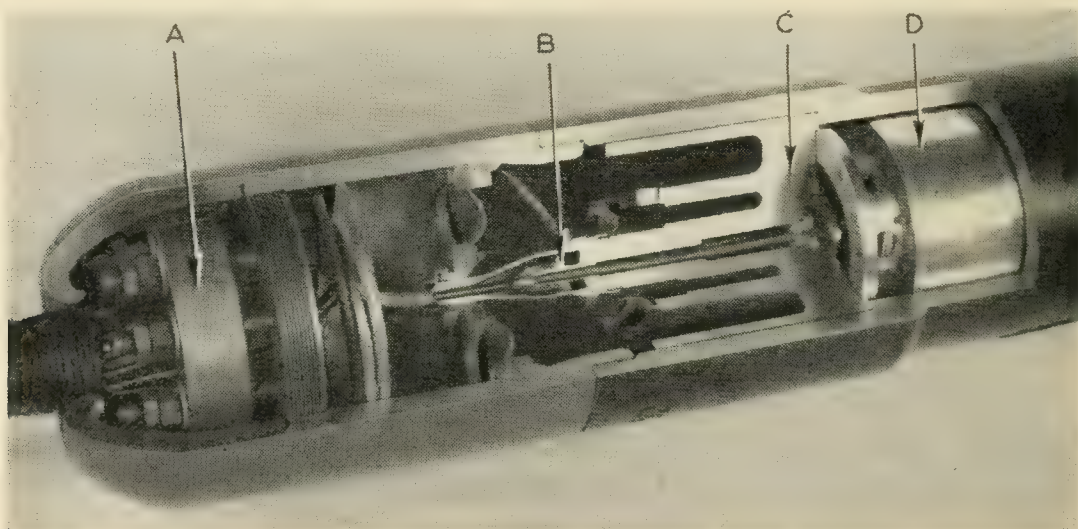


Fig. 8 — Details of housing construction. A. Armour clamp. B. Gland. C. Bulkhead. D. Part of internal unit.

that even with stringent precautions mica is not an entirely satisfactory dielectric material.

Other Components

Electron tubes form the subject of a separate paper.⁷ Of the other miscellaneous components used, one of interest is the short-circuiting fuse across the electron tube heaters. It is constructed like a normal wire-wound resistor on a Sintox tube former, but inside are two cupped copper electrodes filled with a low-melting-point eutectic alloy. If the full line current (316 ma) is passed through the winding, owing to a heater disconnection, the heat generated is sufficient to melt the alloy, which then fuses the two electrodes together. The winding is thus short-circuited, and a permanent connection is left between the electrodes.

DESIGN OF HOUSING AND GLAND

General

Although the maximum depth of water in which British rigid-type repeaters were laid did not exceed about 250 fathoms, the housings used for these repeaters were generally of a type designed for use at ocean depths, and when connected into the cable they were amply strong enough to transmit stresses up to the breaking point of any of the cables used.

The part of the housing which is sealed against water pressure consists essentially of a hollow cylinder, machined from hot-drawn steel tube, and closed at both ends by steel bulkheads carrying the cable glands through which the connections are made to the electrical unit (see Fig. 8). The latter is bolted rigidly to the inner face of the A-end bulkhead. The steel blanks used for the main cylinder and the bulkheads are tested with an ultrasonic crack detector, and after machining they are further subjected to magnetic crack-detection tests.

Each gland has a brass cover which completes the coaxial transmission path and contains a weak solid mixture of polythene and polyisobutylene (p.i.b.). Outside the brass cover is a larger chamber closed by a flexible polyvinylchloride (p.v.c.) diaphragm and containing p.i.b. — a viscous liquid — which prevents sea water coming into direct contact with the bulkhead seal and the gland assembly.

Cylindrical extension pieces, screwed on to the main casing, contain the clamps for attaching the repeater housing to the armour wires of the sea cable, and the housing is completed by dome-shaped end covers. Two external annular ridges near the centre of the housing accommodate

the special quick-release clamp used for handling the repeater during the laying operation.

Protection against corrosion is provided by shot-blasting the surface and then applying hot-sprayed zinc to a thickness of 0.010 in, followed by two coats of vinyl paint. The A-end of the repeater is finished red. The dimensions of the complete repeater are 8 feet $11\frac{3}{8}$ inches $\times$ $10\frac{1}{2}$ inches diameter, and its weight is 1,150 lb in air

Sealing of Housing

The bulkheads, which register on seatings designed to withstand the axial thrust due to the water pressure, are in the form of discs with extended skirts. A watertight and diffusion-proof seal is formed between the casing and the outer skirt of each bulkhead by a silver-soldering process, using carefully controlled electromagnetic induction heating to raise the jointing region to the required temperature. The diametral clearance between the cylinder and the locating surface of the bulkhead is 0.003 inch $\pm$ 0.002 inch, the diameter of the bulkhead being reduced by 0.004 inch for an axial distance of 3 inches from the rim of the skirt to provide a recess into which the molten solder can flow.

The solder is applied as eight pre-formed No. 16 s.w.g. wire rings which are fitted into place cold and coated with a paste formed by mixing flux powder with dehydrated ethyl alcohol. The generator used for heating has a nominal output of 50 kw at a frequency of about 350 kc and is capable of raising the temperature of the jointing region to 750°C in 5 min. The temperature, as indicated by four thermocouples inserted in special holes drilled in the ends of the casing, is maintained at 750°C for a period of 45 min to allow ample time for the entrapped gas and flux pockets to float to the surface. Fig. 9 shows the arrangement for soldering in a bulkhead.

The primary object of the outer skirt is to keep the heated region far enough away from the base to prevent the temperature of the latter rising unduly. Temporary water jackets are also clamped over the gland and around the outside of the casing during the sealing operation. A subsidiary skirt on each bulkhead contains a vent hole which serves to allow displaced air to escape as the bulkheads are inserted into the casing. These vents are later used to apply a low-pressure gas-leak test to the bulkhead seals and then to flush the housing with dry nitrogen to remove any trapped moisture. The vents are finally sealed. At this stage the sealed housing is pressure-tested in water at $1\frac{1}{2}$ tons*/square inch for a

* These are long tons. 1 long ton = 2240 lb.

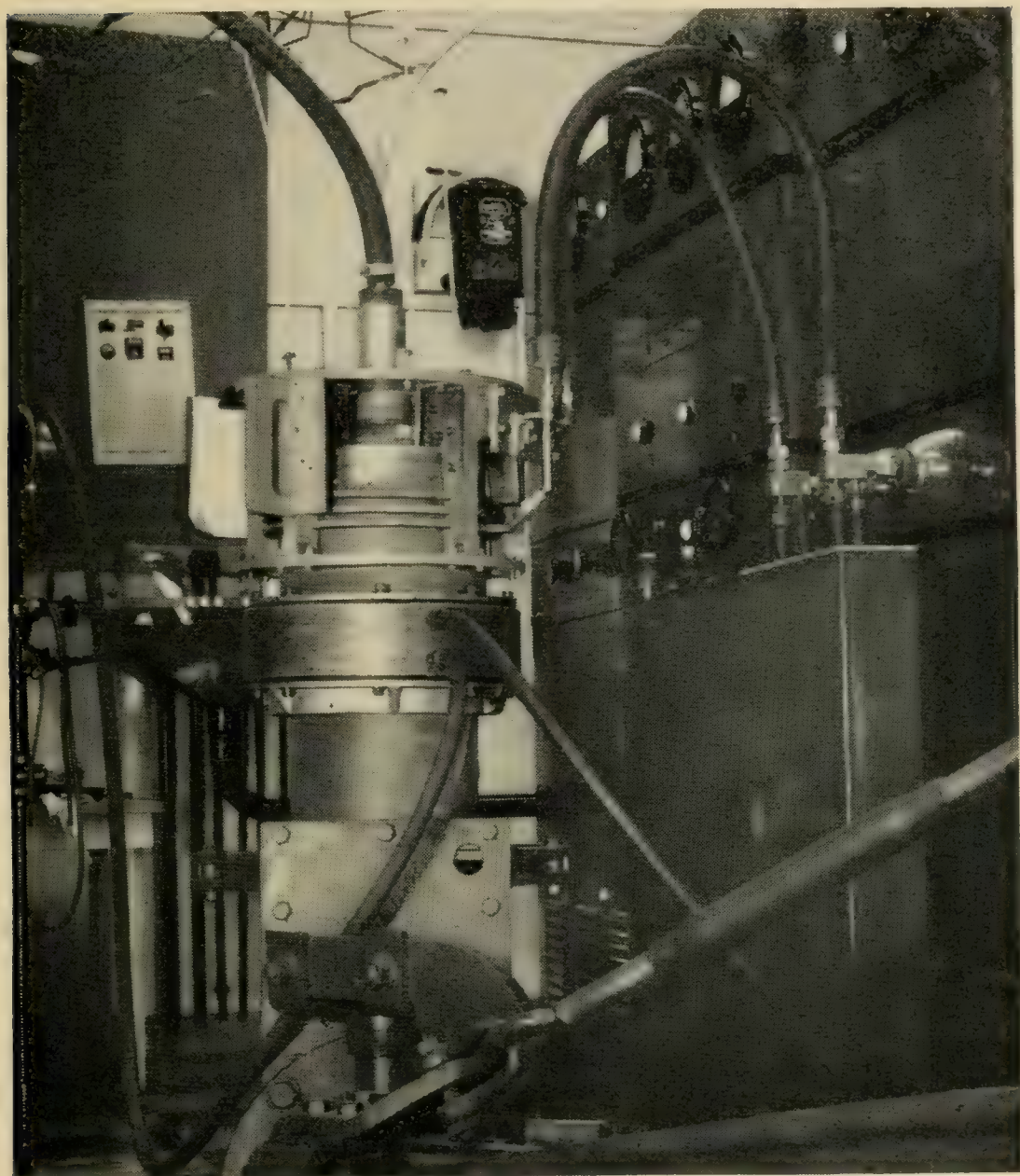


Fig. 9 — Silver-soldering of bulkhead into housing with induction heater.

period of seven days, a moisture detector, mentioned previously, being used to check that no leakage occurs.

Glands

The deep-sea gland was developed from the castellated gland which has been used successfully for a number of years in shallow-water repeaters. The basic principle of this gland is very simple and is shown in Fig. 10; the polythene-insulated cable core passes right through the

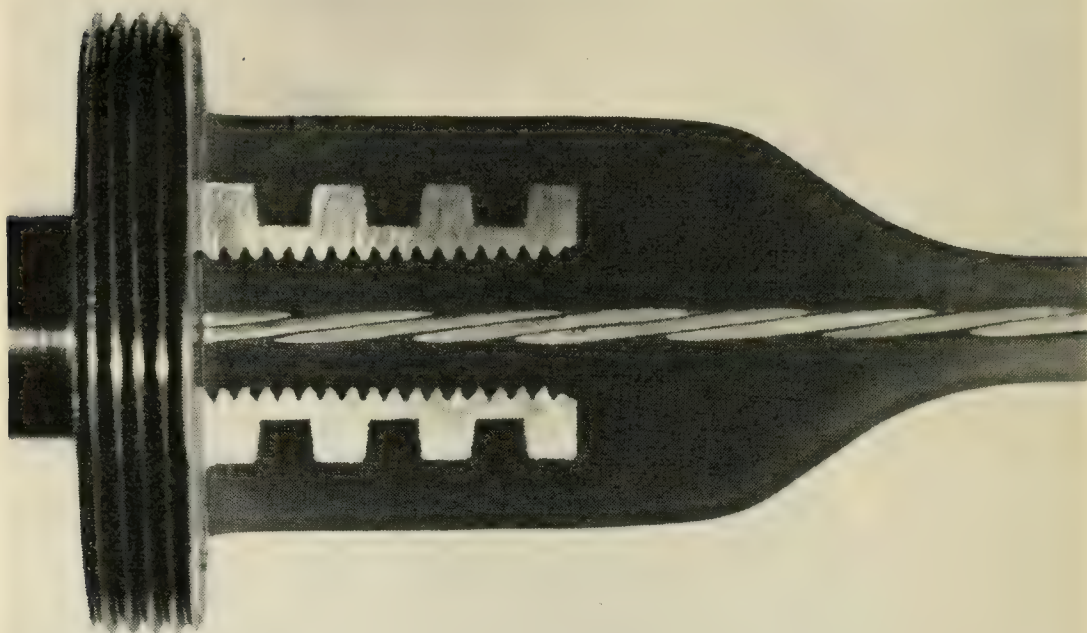


Fig. 10 — Section of high-pressure gland.

bulkhead, and the initial seal is formed by the contraction, during cooling, of polythene moulded on to the core and enclosing a steel stem, having a castellated profile, which forms part of the bulkhead. Each side of the castellation has a taper of about 7° . The application of water pressure increases the contact pressure between the polythene and the stem, thus making the gland inherently self-sealing. As an additional safeguard, the gland stem is first prepared by a lead plating and anodizing process, followed by the application of a thin film of polythene which forms a chemical bond to the plated surface. The injected polythene merges with this film, thus bonding the molded portion to the castellated stem. Tests on sample glands, using a radioactive tracer, have shown no measurable (less than 0.001 mg) water diffusion at a pressure of 5 tons*/square inch over a period of 6 months.

Intrusion of the polythene into the housing, at hydraulic pressures up to at least 6 tons*/square inch, is eliminated by the use of a small-diameter core and by the provision of a screw thread in the hole through which the core passes. During the molding operation a corresponding thread is formed on the polythene core itself. This method of construction distributes the axial force over a sufficiently wide area to prevent any appreciable creep of the polythene.

The completed gland assemblies were all subjected to a minute X-ray examination, followed by a pressure test at 5 tons*/square inch for three months. Whilst under pressure, the glands had to withstand a voltage test of 40 kv dc for 1 minute, to show no ionization effects when a voltage of 3 kv (r.m.s.) at 50 cycles was applied and to have an insulation resistance greater than 20×10^{12} ohms.

MANUFACTURE

General

The manufacture of the electrical units was carried out in accommodation specifically designed for submerged-repeater production. Temperature was controlled at 68°F and the relative humidity was less than 20 per cent in the component shops and 40 per cent in the assembly and test shops. Filtered air forced into the building maintained a slight positive pressure with respect to the outside and eliminated the ingress of dust. With the exception of the tubes and some resistors all components were manufactured in this 'diary' (Fig. 11). Operators are specially selected, and they must change into clean protective clothing in an ante-room before entering the working area, where no smoking or eating is permitted. All operators are particularly encouraged to report or reject any condition which is abnormal or in which they have not complete confidence. Rigorous inspection and testing were carried out by the contractor at all stages, and a Post Office team collaborated with 'floor' inspection and the examination and approval of test results.

The Electrical Unit

The components, after the most careful examination and testing, which in some cases included an aging test, were assembled into their cans and then subjected to a shock test before undergoing detailed electrical characteristics tests. Initial tests on the amplifier were done with a set of 'standard' electron tubes, which were later replaced by the final tubes for the complete tests. The cans were then assembled in a repeater chassis and the electrical tests required before sealing performed — this included a gain response to determine the best settings for the trimmer pads in each transmission band. After fitting and sealing the outer brass case, dry nitrogen was blown through the unit for 24 hours and the gas holes then sealed. The overall electrical characteristics were then taken, the repeater was energized for a two weeks' 'confidence trial' and the characteristics were rechecked. During the 'confidence trial' the gain was

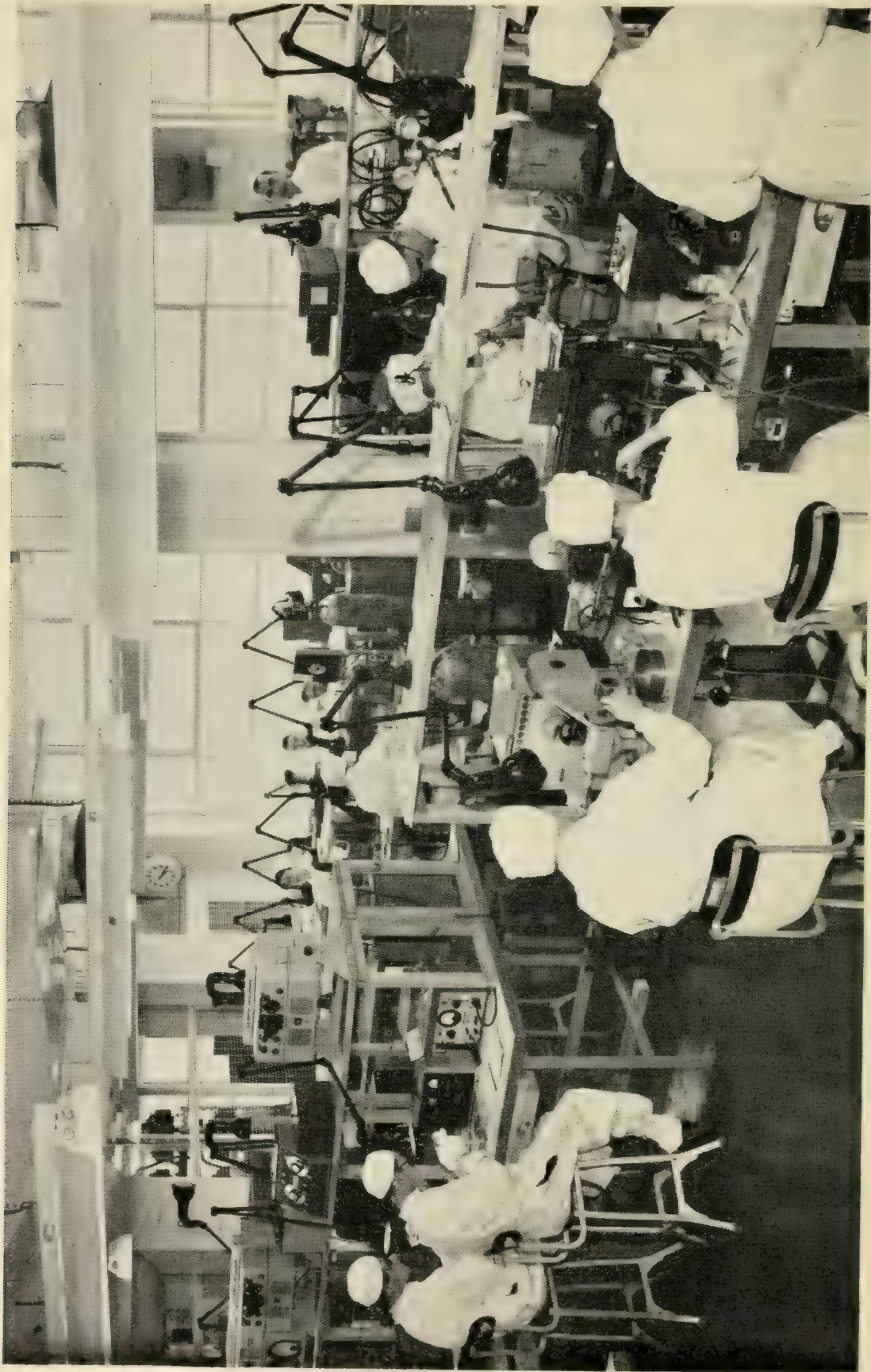


Fig. 11 — Coil production in the 'dairy'.



Fig. 12 — Completed repeater.

continuously monitored on a recorder (duplicated to distinguish between test equipment and repeater variations) on which changes in gain of 0.01 db were clearly indicated. On satisfactory completion of these tests the unit was ready for housing.

Assembly of Electrical Unit in Housing

The first step was to complete the molded joint between the tail cables from the A-end of the electrical unit and the low-pressure side of the appropriate bulkhead. This joint was X-rayed and proof tested at 20 kv dc for 1 minute. The electrical unit was then bolted to the bulkhead, the slack tail cable being correctly coiled into the recess provided, and the whole assembly was lowered into the housing for the first silver-soldering operation. Following this sealing the tail cable joint was made to the B-end bulkhead, which was then lowered into the housing and sealed.

A leak test was then made by applying an internal air pressure of 5 lb/square inch (gauge) and observing the surface when wetted with a solution of a suitable detergent in water. After flushing with dry nitrogen the vents were sealed and the housing was pressure tested. Finally the brass gland covers were fitted and filled with compound, the extension pieces were screwed on, the flexible diaphragms were fitted and the internal space was filled with polyisobutylene. The housing was then ready for further electrical testing.

Tests on Complete Repeater

After housing, the repeaters were submerged in a tank of water for a three-month electrical 'confidence trial.' Before and after the trial the complete characteristics were checked and the noise was monitored on both terminals; during the first and last few weeks of the trial the gain

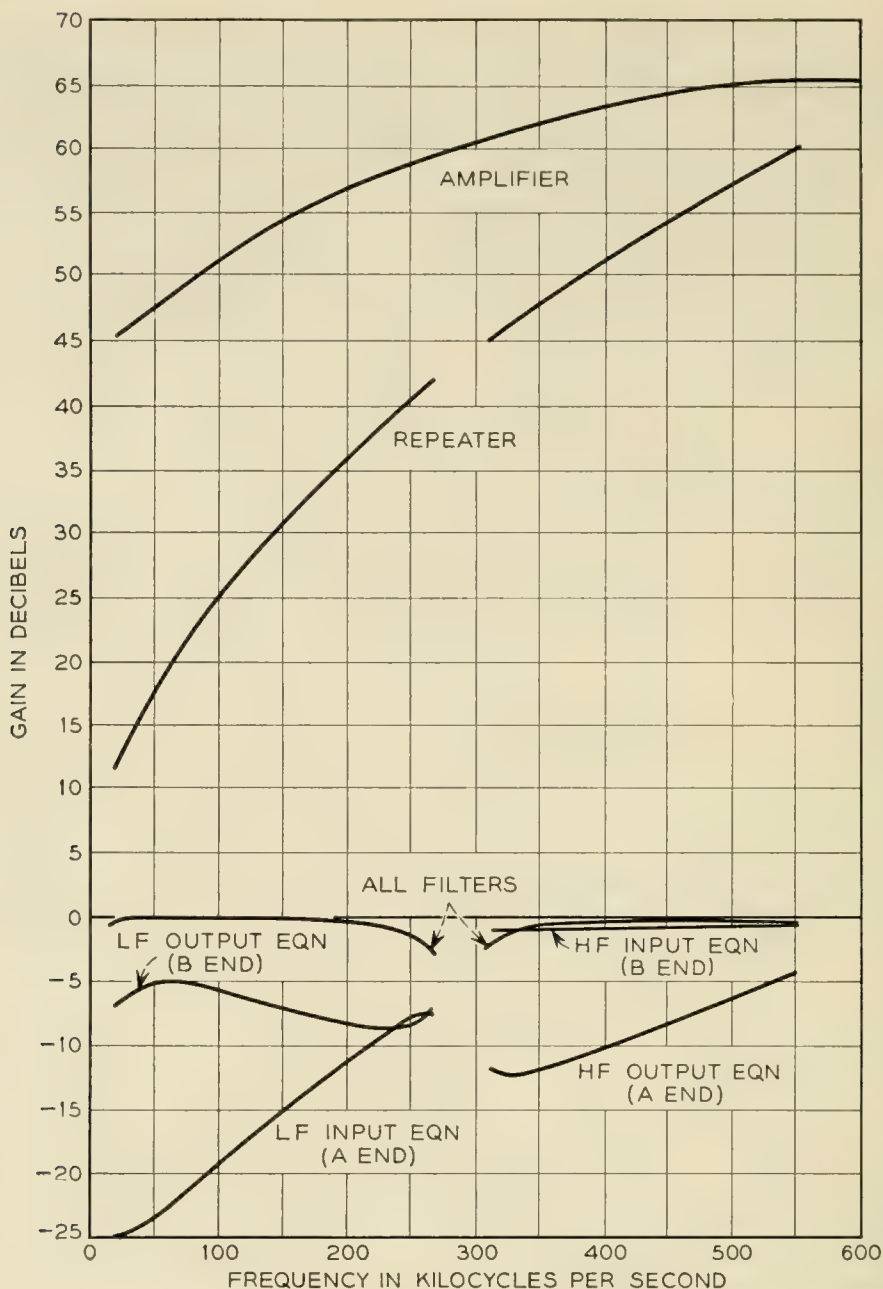


Fig. 13 — Repeater gains and losses.

was monitored on recorders in both directions. Owing to the insertion and withdrawal of repeaters in the power circuit from time to time, the repeaters were subjected to several power-switching operations and temperature cycles.

Stability of electrical characteristics, particularly gain, between the pre-housing tests and the completion of the 'confidence trial' some four months later was regarded as an important criterion of the reliability of a repeater. Unfortunately test conditions and differences between, and

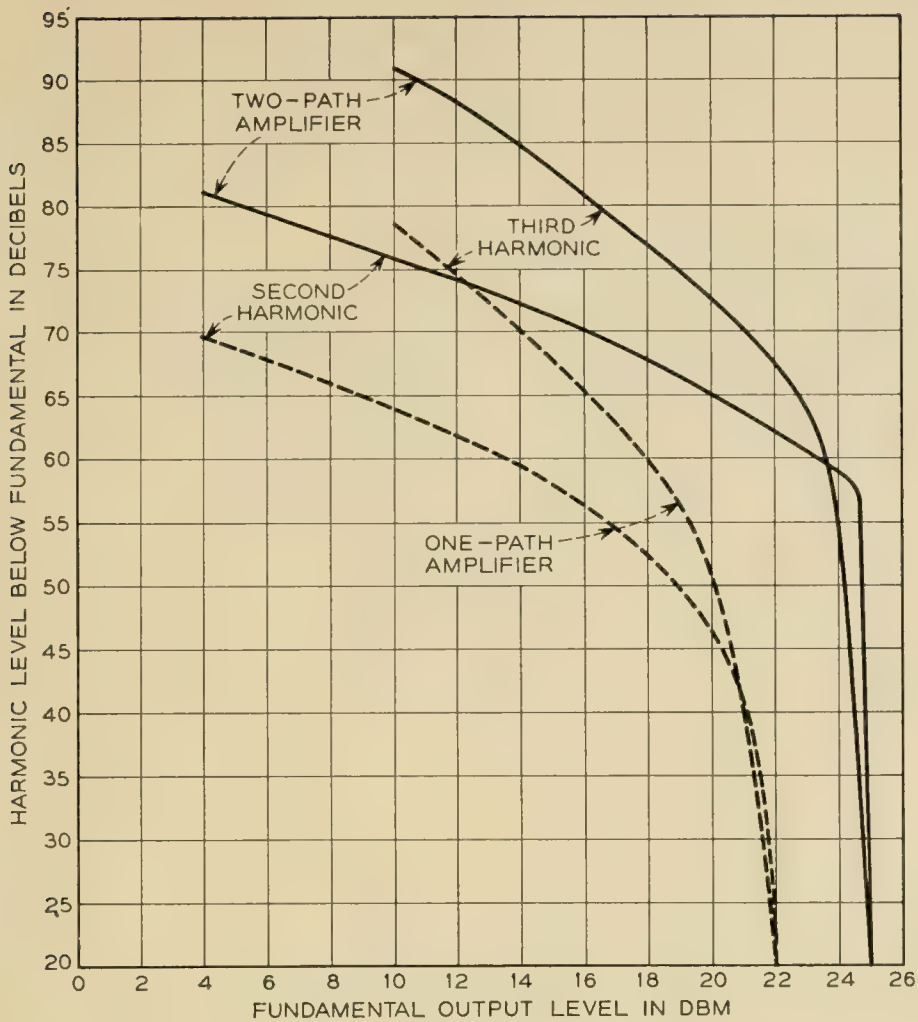


Fig. 14 — Amplifier distortion.

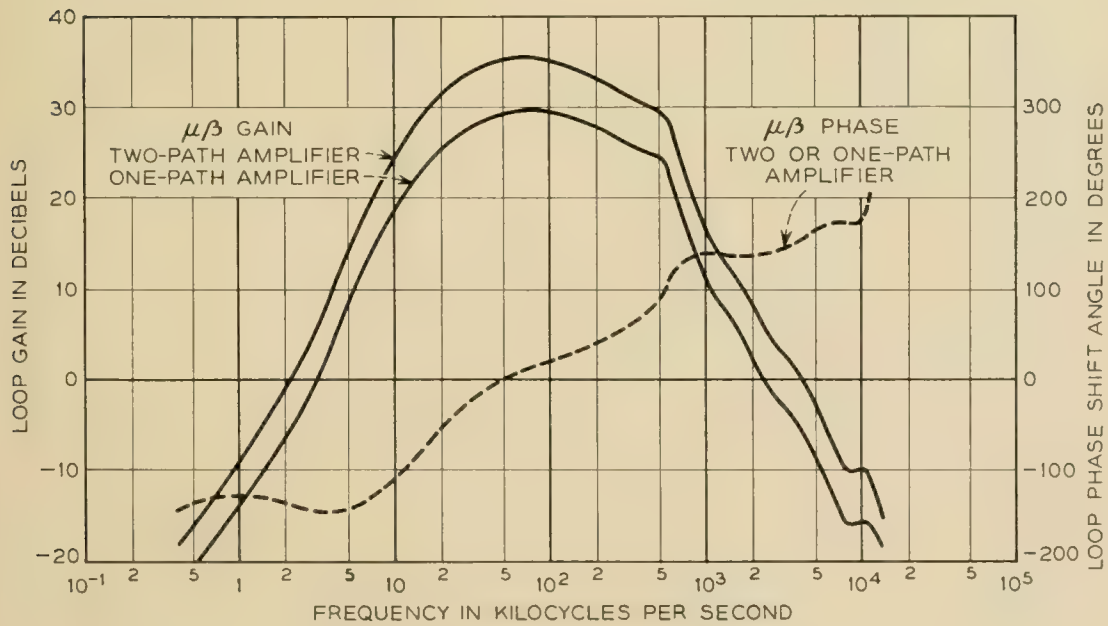


Fig. 15 — Amplifier $\mu\beta$ characteristic.

stability of, the testing equipments reduced the accuracy originally expected, but even so, changes of over 0.1 db were regarded as significant.

Connection of Repeaters to Cable

On board the cable ship the cable ends were prepared by making tapered molded joints to 0.310-inch tail cable, sliding the domed ends of the repeater up the cable and forming the armor wires round the armor clamps. The tapered joint included a castellated ferrule on the center conductor, which, operating on the principle of the main gland, acts as a barrier against the possible passage of water down the center conductor into the repeater. The final assembly operation consists of jointing the tail cables, bolting the armor clamps to the repeater housing and screwing on the domed ends. Fig. 12 shows a completed repeater connected to a cable.

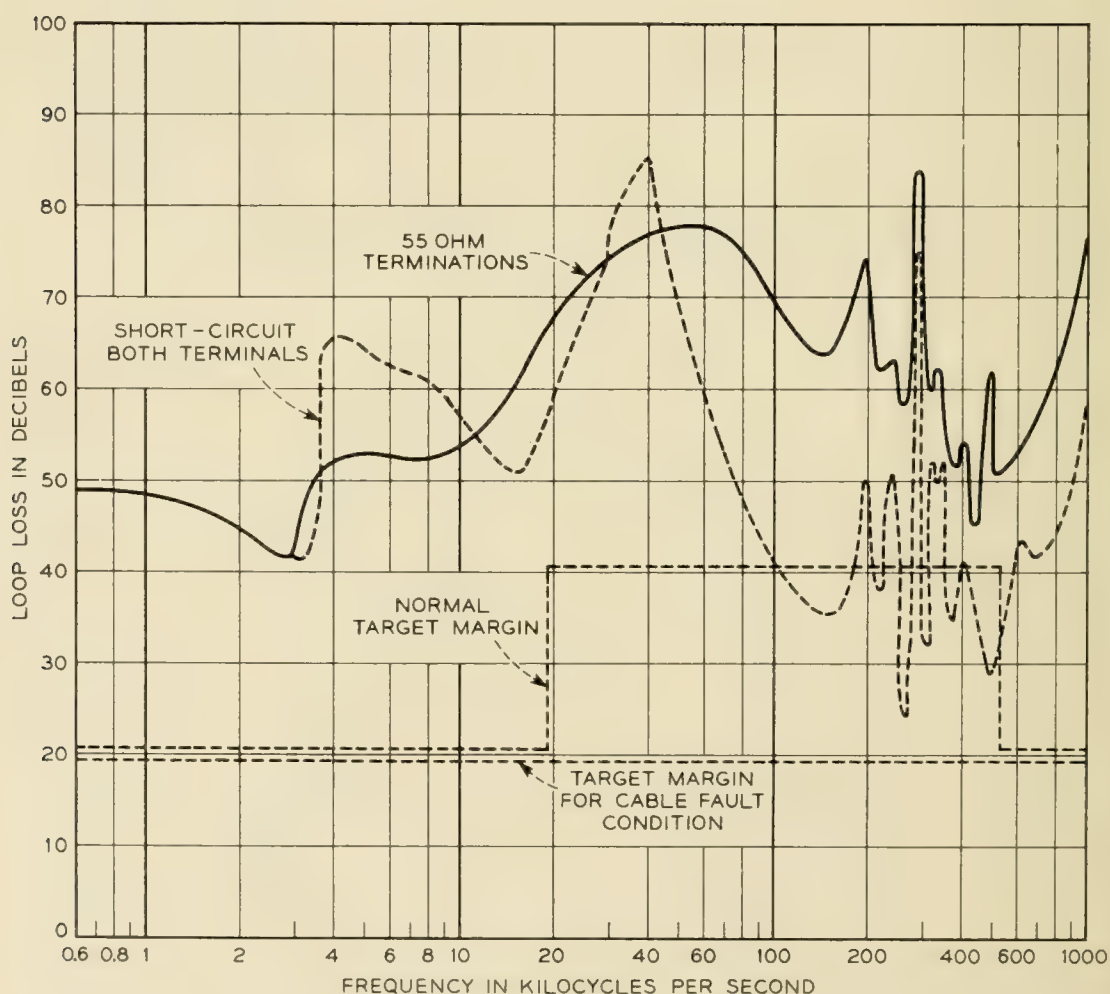


Fig. 16 — Repeater loop loss measured at amplifier input.

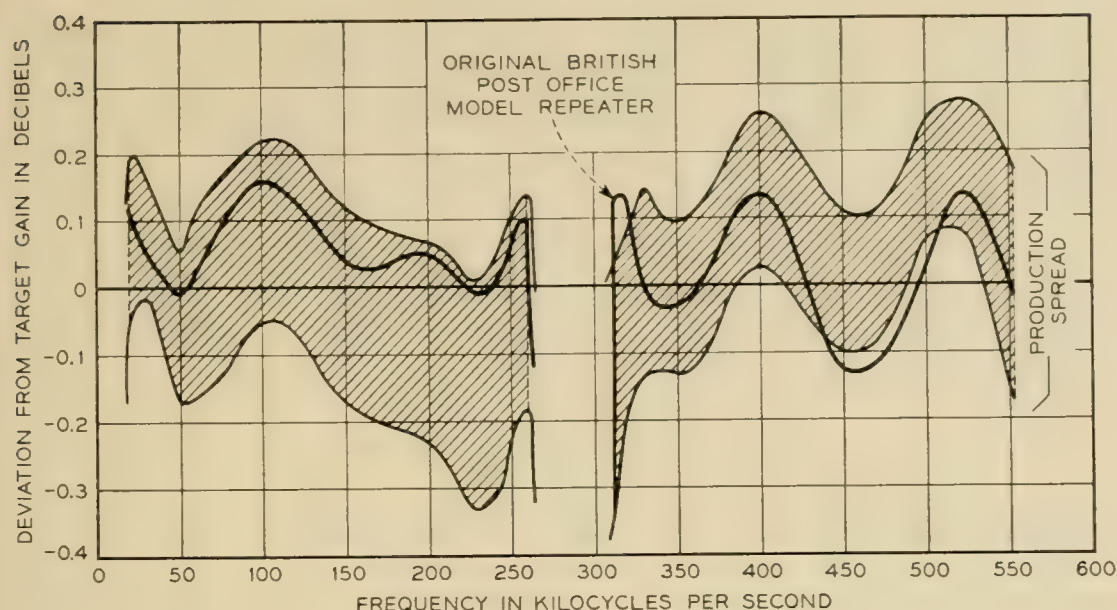


Fig. 17 — Repeater-gain response.

PERFORMANCE

The gains and losses of various sections of the repeater are shown in Fig. 13. Figs. 14 and 15 show, respectively, the harmonic distortion and the stability characteristics of the amplifier with one and two paths operating. The total shunt loss across the amplifier is shown in Fig. 16 as a margin above the amplifier gain. The curves show the result with 55-ohm terminations on the repeater and with a short-circuit on each terminal. Fig. 17 shows the production spread in gain of the 16 repeaters for the system as a deviation from the target value. The highest standard deviation (at 260 kc) was only 0.11 db. In all respects the production repeaters proved to be very consistent and satisfactory in their performance and differed little from the original laboratory-built model.

Typical electrical characteristics of a repeater and the submerged equalizer are shown in Appendices 1 and 2 respectively.

The characteristics of the completed link are described elsewhere,⁹ but it is of interest to note that the overall tests showed that the link behaved as predicted and met the noise requirement and the design margins.

ACKNOWLEDGMENTS

It will be appreciated that the design and manufacture of these repeaters has been an undertaking of teams rather than of individuals. The authors are very grateful to Standard Telephones and Cables, Ltd.,

and Submarine Cables, Ltd., for their unsparing efforts to ensure the very highest standards in the engineering, production and testing of these repeaters. Numerous firms have also co-operated in specialized fields, and the authors are very appreciative of their help. The authors also wish to acknowledge the enthusiastic assistance of their many colleagues in the Research Branch of the Post Office in the design and in the supervisory inspection during manufacture. Finally a grateful acknowledgment is made to the Engineer-in-Chief of the British Post Office for permission to publish the paper.

APPENDICES

1 PERFORMANCE OF TYPICAL SUBMERGED REPEATER

- (a) Insulation resistance 8000 megohms (cold).
- (b) DC resistance at 20°C.

Current, ma	DC resistance, ohms
5	343.0
20	343.2
50	344.2
100	347.9

- (c) Voltage drop at 316 ma 124 volts
- (d) Carrier gain (55 ohms) without moisture detector.

Frequency, kc	Gain, db	Frequency, kc	Gain db
20	11.54	308	44.52
30	13.82	312	44.87
50	17.47	320	45.56
100	25.06	330	46.29
150	30.86	350	47.75
200	35.81	400	51.19
230	38.40	450	54.20
260	40.96	500	57.31
264	41.13	552	60.01
268	40.83		

- (e) Noise level.
 - A terminal (312–552 kc) –59.8 dbm
 - B terminal (20–260 kc) –70.3 dbm
- (f) Harmonic level.
 - 170 kc fundamental level at B terminal +10 dbm
 - 340 kc second harmonic level at A terminal –60 dbm
 - 510 kc third harmonic level at A terminal –58 dbm

(g) Supervisory — 260.800 kc (nominal).

Fundamental level at B terminal, dbm	Second harmonic level at A terminal, dbm
−12	−26
−2	−13.8
+8	−3.6

(h) Moisture detector.

Resonant frequency with 30-ft cable tail..... 1,237 kc

(i) Impedance.

Return loss against 55 ohms

Frequency, kc	A-terminal return loss, db	B-terminal return loss, db
20	17	8
50	17	13
100	16	15
200	16	4
260	16	4
312	13	21
350	14	14
500	18	16
552	16	25

2 PERFORMANCE OF SUBMERGED EQUALIZER

- (a) Insulation resistance..... 8,000 megohms
- (b) DC resistance at 20°C..... 9.2 ohms
- (c) Voltage drop at 316 ma..... 3.0 volts
- (d) Carrier loss (55 ohms) — without moisture detector.

Frequency, kc	Loss, db	Frequency, kc	Loss, db
20	4.28	260	15.90
30	4.23	312	18.26
50	4.87	350	19.96
100	7.47	400	21.75
150	10.35	450	23.16
200	13.04	500	24.55
230	14.52	552	25.97

(e) Moisture detector.

Resonant frequency with 5-ft cable tail..... 1,387 kc

(f) Impedance

Return loss against 55 ohms

Frequency, kc	A-terminal return loss, db	B-terminal return loss, db
20	35	31
50	19	19
100	25	25
260	34	27
552	30	23

REFERENCES

1. R. J. Halsey and F. C. Wright, Submerged Tilybene Repeaters for Shallow Water, Proc. I.E.E., **101**, Part I, p. 167, Feb., 1954.
A. H. Roche and F. O. Roe, The Netherlands-Demark Submerged Repeater System, Proc. I.E.E., **101**, Part I, p. 180, Feb., 1954.
D. C. Walker and J. F. P. Thomas, The British Post Office Standard Submerged Repeater System for Shallow Water Cables (with special mention of the England-Netherlands System), Proc. I.E.E., **101**, Part I, p. 190, Feb., 1954.
2. M. J. Kelly, Sir Gordon Radley, G. W. Gilman and R. J. Halsey, A Transatlantic Telephone Cable, Proc. I.E.E., **102B**, p. 117, Sept., 1954, and Communication and Electronics, **17**, pp. 124-136, March, 1955.
3. J. F. P. Thomas and R. Kelly, Power-Feed System for the Newfoundland-Nova Scotia Link. See page 277 of this issue.
4. B. D. Holbrook and J. T. Dixon, Load Rating Theory for Multi-Channel Amplifiers, B.S.T.J., **18**, p. 624, 1939.
5. R. A. Brookbank and C. A. A. Wass, Non-Linear Distortion in Transmission Systems, J.I.E.E., **92**, Part III, p. 45, 1945.
6. J. S. Jack, Capt. W. H. Leech and H. A. Lewis, Route Selection and Cable Laying for the Transatlantic Cable System. See page 293 of this issue.
7. J. O. McNally, G. H. Metson, E. A. Veazie and M. F. Holmes, Electron Tubes for the Transatlantic Cable System. See page 163 of this issue.
8. S. N. Arnold, Metal Whiskers — A Factor in Design, Proc. 1954 Elec. Comp. Symp., pp. 38-44, May, 1954, and Bell System Monograph 2338.
9. R. J. Halsey and J. F. Bampton, System Design for the Newfoundland-Nova Scotia Link. See page 217 of this issue.

Power-Feed System for the Newfoundland-Nova Scotia Link

By J. F. P. THOMAS* and R. KELLY†

(Manuscript received September 22, 1956)

Design engineers now have available the results of many years of operating experience with submerged-repeater systems supplied from electronic, electromagnetic and rotary-machine power equipments. To meet the very high standards of reliability required for the transatlantic telephone system, a scheme has been evolved that is a combination of new developments and the best features of previous methods. Electronic-electromagnetic equipment forms the basis of an automatic no-break system requiring very little routine maintenance.

INTRODUCTION

The operating power for the submerged repeaters of the Clarenville-Sydney Mines link is derived from a constant current supplied over the central conductor from power equipments located at the two terminal stations. In order to protect the electron tubes in the repeaters the current must be closely maintained at the design value, irrespective of changes in the mains supply voltage or ground potential differences between the two ends of the link. Automatic tripping equipment must be provided to disconnect the cable supply should the current deviate beyond safe limits, but otherwise there must be a minimum of interruptions due to power-equipment and primary-source failures.

Earlier British Post Office schemes have been powered by electronically controlled units feeding from one end only.‡ Manual change-over to standby units has been provided at the end feeding power and at the distant end — a method which has satisfactorily met the economic requirements of short schemes.

* British Post Office, † Standard Telephones and Cable Ltd.

‡ WALKER, D. C., and THOMAS, J. F. P., The British Post Office Standard Submerged-Repeater System for Shallow-Water Cables (with special reference to the England-Netherlands System), *Proc. I.E.E.*, **101**, Part I, p. 190, Feb., 1954.

For the Clarendville-Sydney Mines link a new and more reliable design of equipment has been developed. An automatic no-break system provides an uninterrupted supply to repeaters unless equipments at both ends of the link simultaneously fail to deliver power.

The main improvement in the reliability of the equipment is the replacement of all high-power electron tubes by electromagnetic components. The automatic no-break system takes advantage of the fact that the rating of the repeater isolating capacitors has been chosen to permit single-end feeding. Normally the link is fed from both ends, but in the event of one equipment failing to deliver power the link is powered from the other end without interruption to the cable supply. If the failure is due to a power-equipment fault, double-end feeding can rapidly be re-established by manual switching to the standby. During an ac-supply failure, single-end feeding must be maintained until the supply is restored. In view of the very reliable no-break ac supply provided at both stations, the possibility of a simultaneous ac supply failure at both ends of the link is extremely remote.

The power equipments at each end of the link must be capable of supplying the whole of the power to the cable should one end fail, which requires that each should be capable of operating as a constant-current generator. If two constant-current generators are connected in series, unless precautions are taken, an unstable combination will result and the unit supplying the higher current will drive the other unit 'off load.' Manual adjustment could be provided to equalize the currents fed from the two ends, but a different solution has been developed in which one of the units is a constant-current master and controls the line current, while the other unit is a slave whose voltage/current characteristic in the normal operating range is such that its current is always equal to that of the master unit. If the slave unit fails, the master will take over the supply; if the master unit fails, the slave unit will take over the supply and automatically assume the role of a master unit. The first unit switched on to an unenergized link operates as a master generator and the other unit, on being switched on, automatically operates as a slave. Other than ensuring that the link is safe for energizing, there is no need for any cooperation between the two ends when putting the equipment into service.

DETAILS OF METHOD EMPLOYED

The Master-Slave System of Operations

The output-current/output-voltage characteristics for the equipments are shown in Fig. 1, and are the same for both regular and alternate

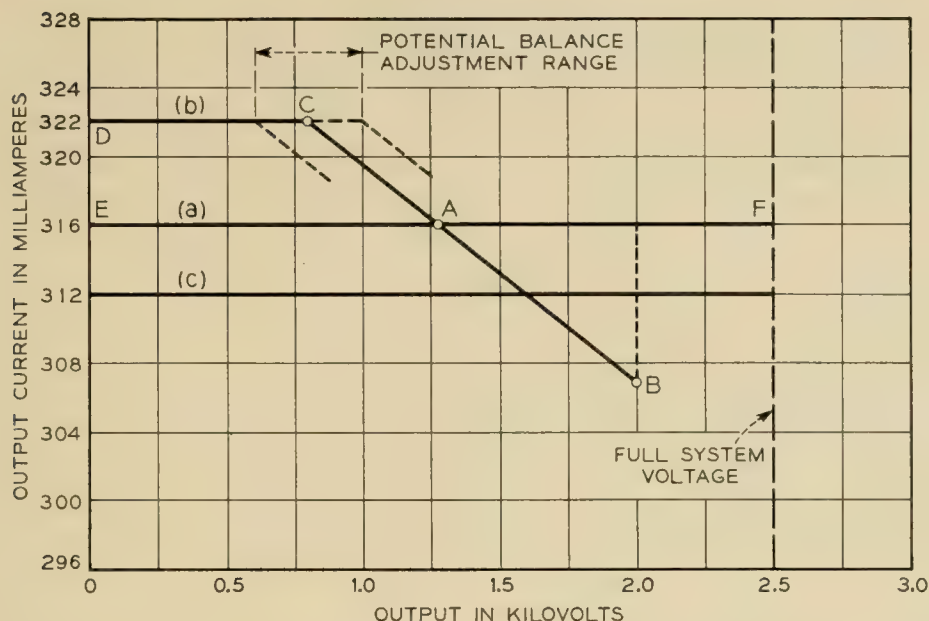


Fig. 1 — Output-current versus output-voltage characteristics. (a) Master. (b) Slave. (c) Master versus master shut-down.

equipments at both ends of the link. Any equipment can operate with any of the characteristics (a), (b) or (c).

Normally the choice between characteristics (a) and (b) is made automatically by the equipment. If the output voltage does not reach 80 per cent of the full link voltage (approximately 2 kv), the unit will have the slave characteristic (b) (DCAB). If the output voltage reaches or exceeds 2 kv, the unit automatically switches to the master characteristic (a) (EAF). Once having switched to the master characteristic the equipment does not automatically change back to the slave characteristic even if the output voltage falls below 2 kv.

When the link is energized, the first equipment switched to line will come on as a slave unit; its output voltage will then pass 2 kv and it will automatically be switched to the master characteristics and the complete link will be energized from one end. The second unit switched to line will come on as a slave unit, and having a higher output current, will drive down the output voltage of the master unit at the other end until the current of the slave unit has become equal to that of the master (DCA in Fig. 1). The output voltage of this equipment will not exceed 2 kv and it will not switch to master. The cross-over point of the two characteristics, A, will determine the potential fed from each end, and this can be adjusted as indicated by the dotted lines near C.

In an emergency, an equipment can be taken out of service by switching off the ac supply, and the links will then be powered from the dis-

tant end only. Should the master end be switched off, the distant slave unit will switch to master characteristic as soon as its output voltage exceeds 2 kv. This abrupt disconnection of one equipment causes unnecessary voltage surges on the cable, and when an equipment is removed for normal maintenance purposes the following procedure is adopted. The slave equipment is switched to the master characteristic (an external key is provided for changing from master to slave or vice versa, but see the limitation described in the next paragraph). This leaves two master equipments feeding the cable, but any redistribution of voltage is slow since the currents are approximately equal. An external control (master/master shut-down), is then operated on the equipment to be taken out of service, changing its characteristic to that shown in Fig. 1(c). The output voltage of this unit will then be slowly reduced, and when it is zero the equipment can be switched off without causing surges on the cable.

The current deviations occurring over the slave characteristic (c) (from +2 per cent to -3 per cent) are the maximum permitted by the tube design engineers for the tubes in the submerged repeaters, and are permitted only for short periods. The circuit associated with the manual switching of the equipment to the slave characteristic is therefore made inoperative if the output voltage exceeds 2 kv, since in this range the slave characteristic is the extension of the line CAB on curve (b) and the current is outside the permitted range. The slave characteris-

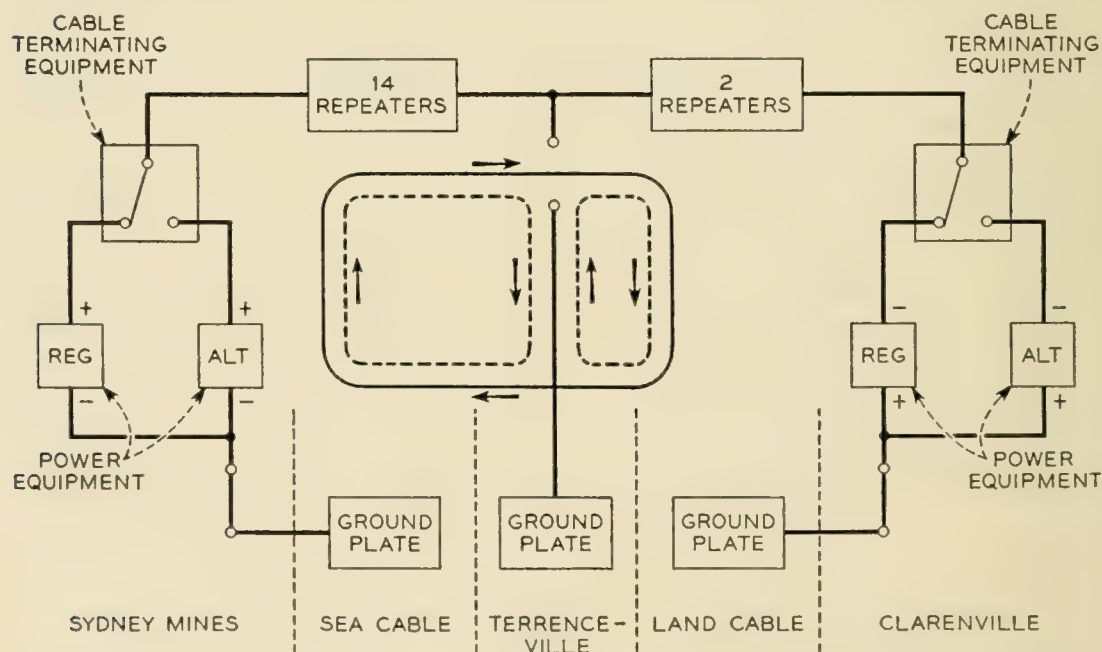


Fig. 2 — Current paths on Clarendville-Sydney Mines link.

tic over the range CAB is controlled by a voltage-sensitive circuit connected near the output of the equipment, and the stability of the characteristic against input voltage and component aging is the same as for the master characteristic. Changes in the distribution of the system potential due to supply variations and component aging are therefore small.

Overall Current and Voltage Distribution

Facilities have been provided at the junction of the land and sea cables (Terrenceville) to connect a power ground to the center conductor of the cable. Normally this ground will be disconnected, but during installation it enables the land and sea sections to be energized separately and may subsequently be of assistance in the localization of cable faults near Terrenceville.

The full line in Fig. 2 shows the current path with normal double-end feeding, while the broken lines show the current paths when a ground is connected at Terrenceville.

The full line (d) in Fig. 3 shows the voltage distribution along the link

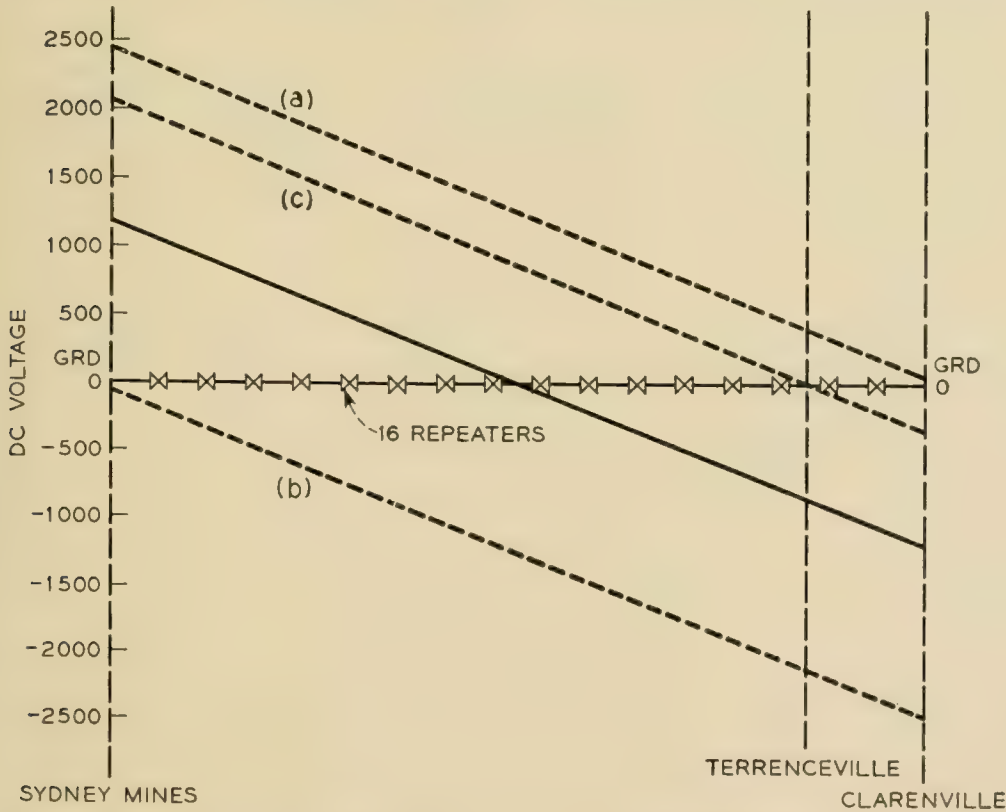


Fig. 3 — Potential distribution on Clarenville-Sydney Mines link. (a) Single-end feeding from Sydney Mines. (b) Single-end feeding from Clarenville. (c) Ground at Terrenceville; double-end feeding. (d) Normal operation; double-end feeding.

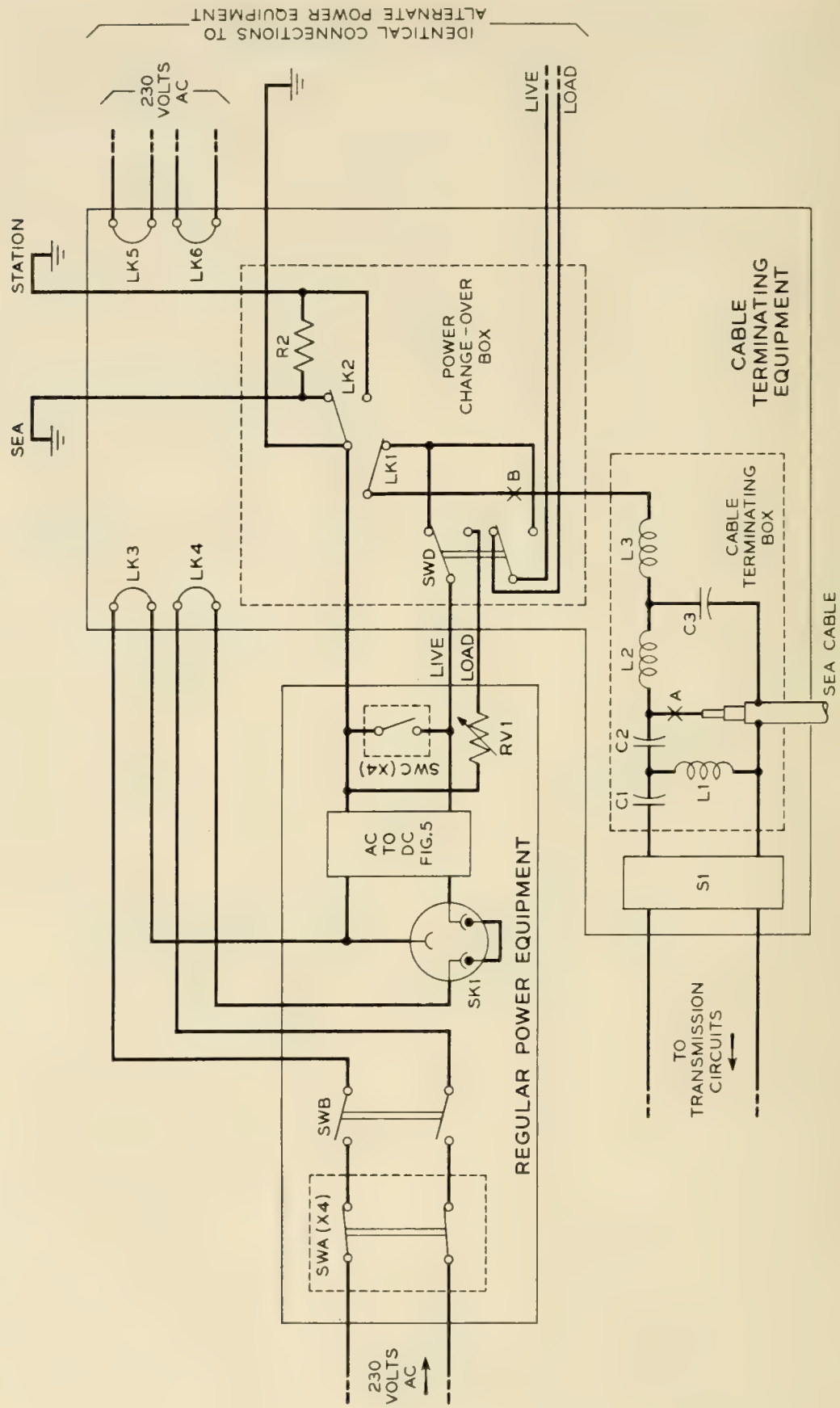


Fig. 4 — Cable terminating equipment and interconnections with power equipments.

with normal double-end feeding, the broken lines (a) and (b) show the distribution with single-end feeding from Sydney Mines and Clarenville respectively and the broken line (c) shows the distribution when a power ground is connected at Terrenceville.

DETAILS OF EQUIPMENTS

Connection of the Equipment to the Cable

Each station is provided with two power equipments and one cable-terminating equipment interconnected as shown in Fig. 4. When the live side of the output of the regular equipment is connected to the cable via SWD and the link LK1, the alternate equipment output is connected to its own dummy load (equivalent of RV1) and vice versa.

The grounded sides of the regular and alternate equipments are made common and then connected via a removable link, LK2, to the sea ground. A safety resistor R2 connects the sea ground to the station ground to restrict the rise in potential to 100 volts if the sea ground becomes disconnected. During maintenance on the sea-ground circuit the link LK2 can connect the power equipment ground to the station ground.

If, for maintenance purposes, it is necessary to feed from the distant end only, the link LK1 can connect the cable to the power-equipment ground and disconnect the live side of both power equipments from the line.

Safety Interlocks

Safety interlock circuits are installed to protect the maintenance staff if the equipments are used incorrectly and are not the normal methods employed for controlling the power supplies. With double-end feeding, dangerous voltages are generated at both ends of the system, and when access is gained to any point in the equipments personnel must be protected from the local and distant power sources.

To minimize interruptions to traffic, the units of the cable-terminating equipment have been grouped under three headings (see Fig. 4), namely

(a) Transmission equipment (S1) isolated from the dc cable supply: this includes cable simulators and monitoring facilities not associated with the power supplies.

(b) Equipment associated with the dc supply that can be made safe without interrupting traffic.

(c) Equipment that can be made safe only by interrupting traffic.

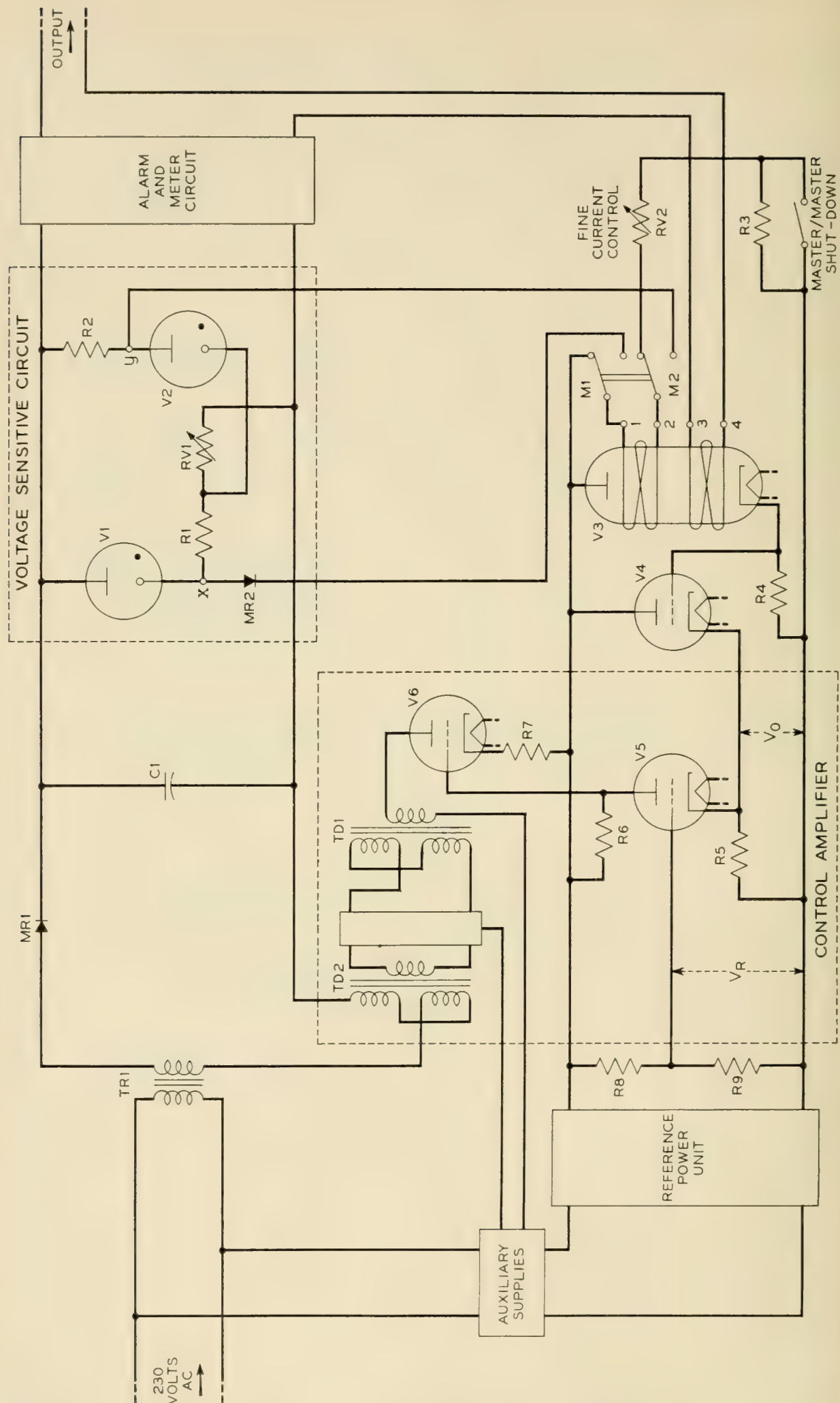


Fig. 5 — Simplified schematic of power-feeding equipment.

Group (a) is treated as normal terminal transmission equipment and is not interlocked. Groups (b) and (c), the power-change-over and cable-terminating boxes respectively, both require that the local power equipments are switched off before they can be made safe. The external panel covers over both boxes have links disconnecting the ac supplies to both of the local power equipments when the covers are removed (LK3-LK6). This will not interrupt traffic, since single-end feeding will continue from the distant station.

The doorknob switch of the power-change-over box automatically connects the live lead to ground (at the point B) when the door is opened. This will not interrupt traffic since the impedance of L2 is high at carrier frequencies. Changing over the power equipments and grounding arrangements can therefore be performed without interrupting traffic.

The doorknob switch of the cable-terminating box automatically connects the center conductor of the cable to ground (at the point A) when the door is opened. Traffic will therefore be interrupted if maintenance work is necessary within this box.

Access is gained to a power equipment by opening one or more of four doors. Each of the doorknob switches disconnects the ac supply (SWA) and short-circuits the output of the equipment (SWC).

It will be appreciated that if the correct procedure is adopted the local power equipments will be switched off by SWB after the master/master shut-down procedure (described in section on *The Master-Slave Systems of Operations*) has been carried out and not by removing the panel covers or opening the doors.

Electrical Details of Power Equipment

Control Method.

Fig. 5 is a functional simplified circuit diagram in which TR1, MR1 and C1 represent a conventional unregulated power unit.

The control circuit compares a signal proportional to the output current, V_o , with a stable reference signal, V_R , the two signals being applied to the cathode and grid, respectively, of V5; the difference between them is amplified and used to adjust the voltage fed to the rectifier circuit, MR1, to maintain the output current constant.

Ignoring at this stage the auxiliary coil 1-2 of the magnetically controlled diode V3, the output current flows through the coil 3-4 and controls the voltage across R4 and hence, via the cathode-follower V4, the potential V_o on the cathode of V5. V_R is a function of the output of a conventional electronic constant-voltage power unit (reference power

unit) using a neon tube for its reference. The grid-cathode bias of V5 controls the secondary impedance of the transducer TD2 via the amplifier V5, V6, TD1. The gain of the control amplifier and the sign and magnitude of the normal impedance of the secondary of TD2 are arranged to give a fall of approximately 0.25 per cent in the current from full system load to short-circuit.

Two advantages are obtained by employing a magnetically controlled diode for V3 instead of an electrostatically controlled tube. The control is directly proportional to the output current, and is not dependent upon the stability of a series resistor, while the control circuits are isolated from the output-circuit voltage, which simplifies maintenance of the more complex parts of the equipment.

Current Characteristics.

Without current flowing in the auxiliary coil 1-2 (on V3 in Fig. 5) the equipment has a normal constant-current characteristic, the value of which, 322 ma, is preset by adjusting the mechanical position of the coil assembly along the main axis of V3.

Two neon tubes and two resistors form the bridge V1, V2, R1 and R2, which is balanced when the voltage drop across R1 and R2 equals the constant voltage across V2 and V1, the output voltage at which this occurs being adjusted by RV1. At voltages below balance, current tries to flow from y to x and at voltages above balance it flows from x to y. The balance voltage corresponds to the voltage at which the slave-unit characteristic changes slope (C in Fig. 1). For voltages below balance the rectifier MR2 prevents current from flowing in the winding 1-2 (on V3), and in this range the constant current of 322 ma is maintained (see DC, Fig. 1). For voltages above balance, current flows in the winding 1-2 and progressively decreases the output current (CAB, Fig. 1). At 80 per cent of the full link voltage, the contacts of relay M are operated and disconnect the auxiliary coil 1-2 from the voltage-sensitive bridge and connect it across the reference supply. The current through 1-2 is then set by the fine current control (RV2) to make the output current the required 316 ma. As previously stated, relay M does not automatically switch back when the output voltage drops below 80 per cent of full link voltage; the current characteristic of a master unit is EAF in Fig. 1.

The master/master shut-down characteristic [Fig. 1, curve (c)] is obtained by short-circuiting R3, which causes the constant current to fall to 312 ma. Adjusting RV2 would be equally effective, but short-circuiting R3 does not permanently disturb the normal current setting.

Alarms.

The equipment trips and gives both aural and visual alarms for +20 per cent current, +20 per cent full link voltage and for the failure of certain auxiliary supplies that would damage the equipment.

Aural and visual alarms are provided for ± 1 per cent current and ± 3 per cent voltage and for equipment changes from slave to master characteristic or vice versa. Visual indication is given if the ac supply fails.

Reliability.

Where possible, only components of proven integrity have been used. High-power thermionic tubes have been excluded and the tube types employed have been specially selected for long life. Electrolytic capacitors have been excluded from all except one position, and in this case the component has been divided into six units in parallel, the failure of all but one of these units causing only a slight increase in the output ripple.

Particular attention has been given to the continuity of the output circuit. A failure of a power equipment for any reason other than an open-circuit in the output will not interrupt traffic, the link changing to single-end feeding from the far end. A disconnection anywhere in the dc feed path will disconnect all power from the line. Relatively short-lived components in this part of the circuit have either been duplicated in parallel or shunted by resistors capable of carrying the full line current.

Other facilities.

Each equipment has facilities for checking its overall performance. A variable-ratio transformer can be introduced at SK1 (Fig. 4) and with the 4-position dummy load referred to earlier, the regulation against alternating input voltage and output load can be measured.

Provision is made for checking all the alarms, and the current can be measured at strategic points in the control circuit either when the equipment is normal or with the loop feedback disconnected.

Separate large-size meters are provided for measuring the output voltage and current to an accuracy of ± 1 per cent. A more accurate measurement of current is obtained from potentiometric measurements made across a standard resistor connected in series with the output.

Electrical Details of Cable-Terminating Equipment

The majority of the electrical features of the cable-terminating equipment have already been considered.

Within the cable-terminating box (Fig. 4) are the power-separating filters; the high-pass filters C1, C2 and L1 passing the carrier frequencies and the low-pass filter L2 passing the direct current. The transmission equipment represented by the block S1 is for convenience mounted in the cable-terminating cubicle. The extra low-pass filter L3, C3 can be specially designed to prevent signals that are peculiar to the site (local radio stations, etc.), which are picked up in the power equipment, from being fed to the transmission circuits.

Metering facilities (not shown in Fig. 4) are provided to check the continuity of the sea-ground circuit. A separate insulated wire is connected to the sea-ground plate and the continuity is checked by measuring the voltage drop along the ground cable. Aural and visual alarms are given if the sea ground becomes disconnected.

Current and Voltage Recorders

Current and voltage recorders are provided for both the regular and the alternate equipments, the values being measured at the points where the outputs of the power equipments enter the cable-terminating equipment.

A magnetic-amplifier unit drives the current recorder, the scale deflection being from -5 per cent to $+5$ per cent of the normal line current. The voltage recorder is connected across the ground end of a resistance potentiometer, the full-scale deflection being 3 kv.

The magnetic amplifier and potentiometer units are mounted in the cable-terminating equipment, but the four recorders and their associated supplies and alarms are mounted on a separate rack.

Mechanical Details

Fig. 6 shows two power equipments and a cable-terminating equipment as installed at each terminal station. The same cubicle frameworks are used for power and cable-terminating equipments, the power equipment consisting of two cubicles bolted side by side and fitted with doors at the front and back.

The top of the cable-terminating cubicle contains the power-change-over box, the center contains the cable-terminating box, while the transmission and ground-cable test circuits are located near the bottom. The

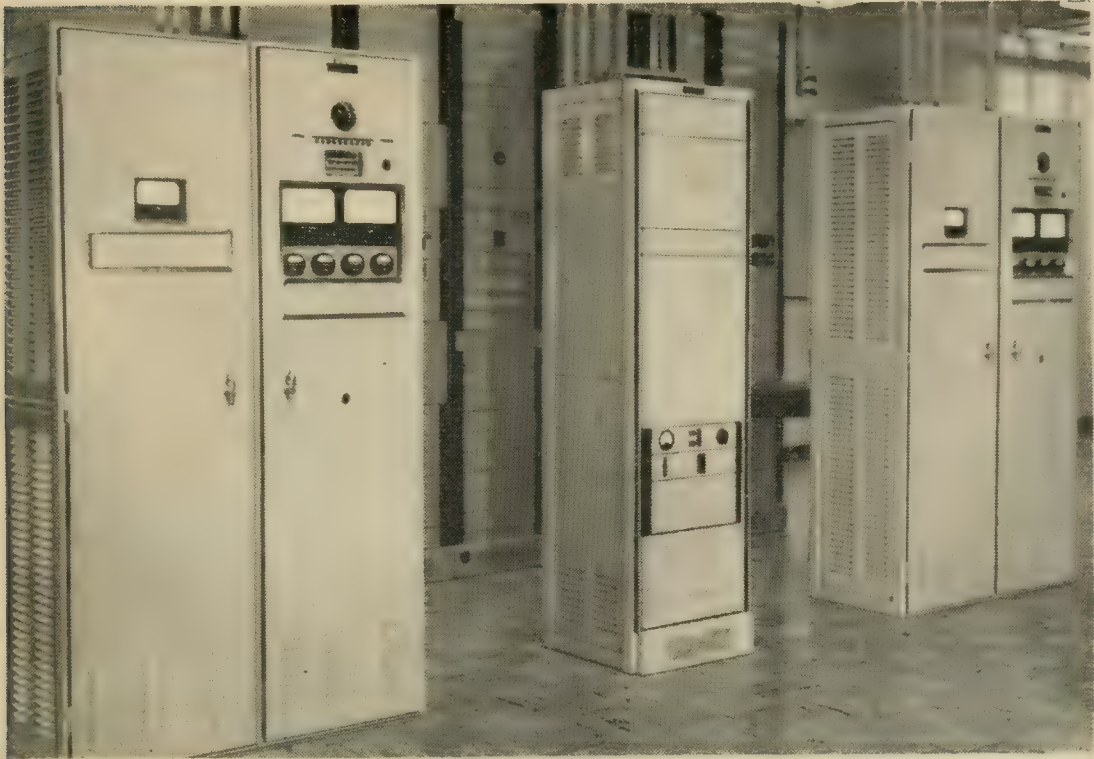


Fig. 6 — Power and cable-terminating equipments as installed at each terminal station.

main cable enters the bay at the top and passes behind the power-change-over box into the cable-terminating box. Access to the recorder units and cable simulators is from the rear.

Fig. 7 shows the front of one power equipment with the doors open. The left-hand cubicle contains the main h.v. transformer, the electronic and magnetic parts of the control circuits, the reference power unit and the auxiliary supplies. The other cubicle contains the main rectifier and smoothing circuits, and the associated meter and alarm circuits.

PERFORMANCE

Six power equipments and four cable-terminating equipments were manufactured for the link. Of these, four power equipments and two cable-terminating equipments were provided for the terminal stations and the remainder were for use on H.M.T.S. *Monarch* during the cable laying and subsequently as off-station and training spares.

After individual testing, the power equipments were checked in pairs energizing a 16-section artificial cable constructed to simulate the Clarenville-Sydney Mines link. These tests were kept running for approximately two weeks on each equipment, and the line current was monitored with

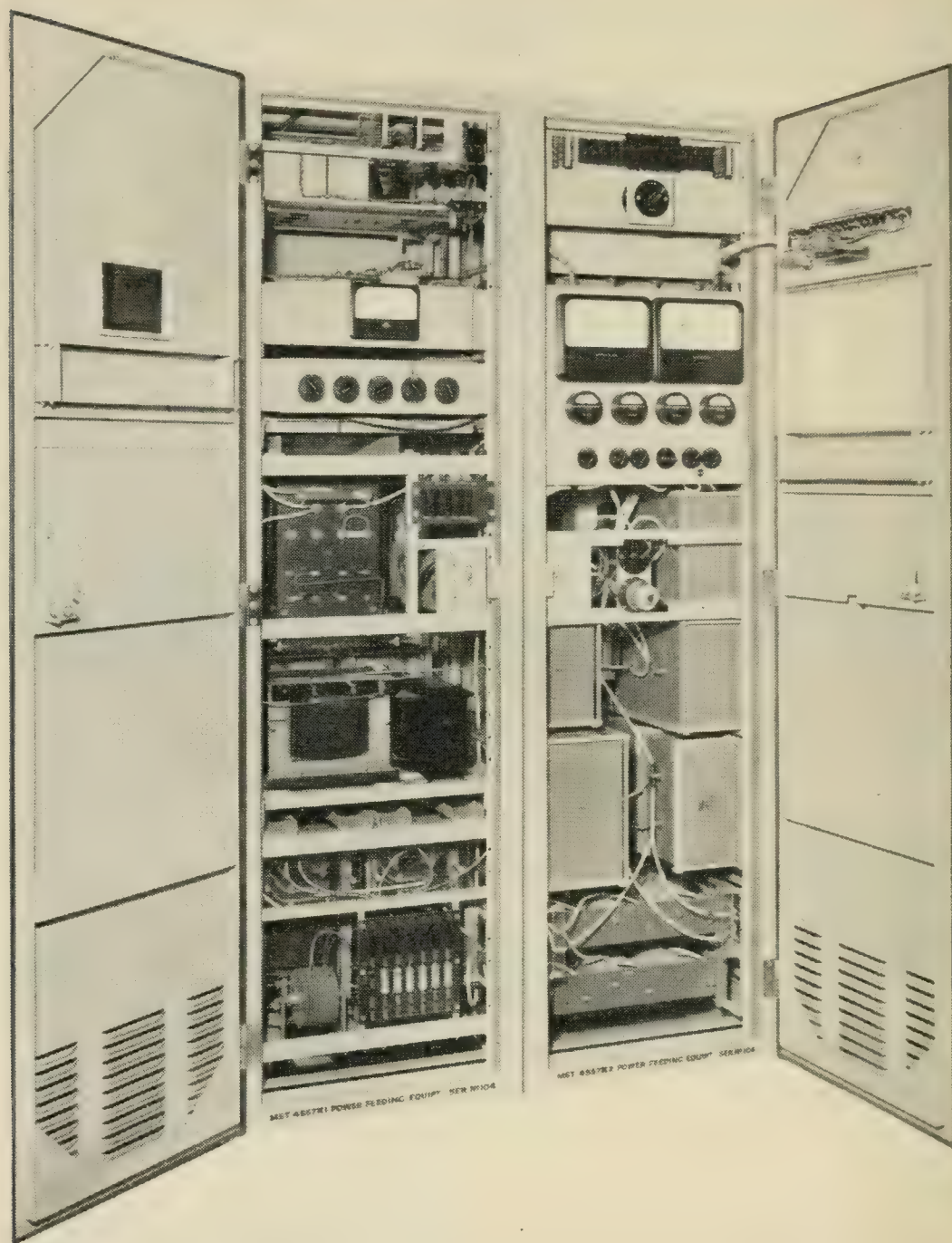


Fig. 7 — Power equipment.

expanded-scale current recorders (from -1 per cent to $+1$ per cent of the normal current).

The equipments were completed at the beginning of 1956, and have given satisfactory service both at the terminal stations and on board H.M.T.S. *Monarch*. From March to June, 1956, all four units at the

terminal stations were operating into their dummy loads and daily current and voltage readings were well within the required specifications. Since the laying of the sea section the equipments have satisfactorily operated the completed link. When the link is in service the working equipment at each end will be changed at 6-monthly periods and the change-over time will be staggered by three months each end. The equipment coming out of service will be immediately routine checked and adjusted if necessary.

The tests since installation confirm the laboratory and factory results that the regulation is better than ± 1 ma for any alternating input voltage from 195–265 volts in the frequency range 40–70 cycles for output voltages of 0 to 3 kv. The relatively long correction period of the magnetic-amplifier control circuit (about 0.5 sec) is satisfactory with the type of no-break supply installed at the stations.

Another useful by-product of double-end feeding is that, when there is a shunt fault on the system, the voltages supplied from the two terminal stations give an indication of the fault position. Many factors will control the accuracy of the location, e.g. magnitude and position of the fault, and how nearly the currents fed from the two ends are equal. Calculations, confirmed by tests made with shunt faults introduced at Terrenceville, show that any continuous shunt fault that will affect transmission (to the extent of operating the 1 db pilot alarms) will be detected to an accuracy of ± 1 per cent. The limit is set by the accuracy with which the link voltage distribution can be measured.

CONCLUSIONS

The power-feeding system described is suitable for submerged-repeater schemes that can, in an emergency, be temporarily powered from one end only. For locations within reasonable reach of a central catastrophe-spare store the equipment should be sufficiently reliable not to need duplication at both terminal stations. For future schemes this will provide a method which is economically attractive compared with the present single-ended methods and which offers numerous electrical advantages.

The routine maintenance required on the power equipment could be further reduced if the few remaining electron tubes were replaced by electromagnetic components. On schemes where short interruptions to traffic do not involve a relatively high loss of revenue it would then be unnecessary for the local staff to maintain the high-voltage equipment and the expensive no-break ac supplies could be abandoned. The latter

depends upon the probability of the primary sources at the terminal stations failing simultaneously.

ACKNOWLEDGMENTS.

The authors wish to express their thanks to many of their colleagues in the Post Office and Standard Telephones and Cables, Ltd., who have contributed to the developments described in the paper. The permission of the Engineer-in-Chief of the Post Office, and of Standard Telephones and Cables, Ltd., to make use of the information contained in the paper is also gratefully acknowledged. The equipments for the Clarendville-Sydney Mines link were manufactured by Standard Telephones and Cables, Ltd., at North Woolwich, London.

Route Selection and Cable Laying for the Transatlantic Cable System

By J. S. JACK*, CAPT. W. H. LEECH† and H. A. LEWIS‡

(Manuscript received September 7, 1956)

The repeatered submarine cables which form the backbone of the transatlantic telephone cable project were installed during the good weather periods of 1955 and 1956. This paper considers the factors entering into the selection of the routes, describes the planning and execution of the laying task and presents a few observations on the human side of the venture. It also covers briefly the routing of some 55 nautical miles of repeatered submarine type cable which were trenched in across the neck of the Burin Peninsula in Newfoundland to connect the Terrenceville submarine terminus with the cable station at Clarenville.

GENERAL

In the days of Cyrus Field, Lord Kelvin and those other foresighted and courageous entrepreneurs of the early transoceanic submarine cable era, the risks involved in selecting a route and laying such a cable must have appeared formidable beyond description. And indeed they were, for not until the third attempt was a cable successfully laid.

Today the hazards may be somewhat more predictable, our knowledge of the ocean bottom more refined and our tools improved, but only to a degree. The task still remains extremely exacting in its demands for sound engineering judgment, careful preparation, high grade seamanship, and good luck — weatherwise. For the basic methods now in use are still remarkably like those employed on *Great Eastern* and other early cable ships and the meteorological, geographical and topographical problems have changed not at all.

In the current transatlantic project — the first transoceanic telephone cable system — there are two submarine links. Between Clarenville in Newfoundland, and Oban in Scotland there lie some 1,850 nautical miles§ of North Atlantic water, most of it deep and all of it subject to

* American Telephone and Telegraph Company. † British Post Office. ‡ Bell Telephone Laboratories.

§ A nautical mile, as used herein, is 6,087 feet, 15.3 per cent longer than a statute mile.

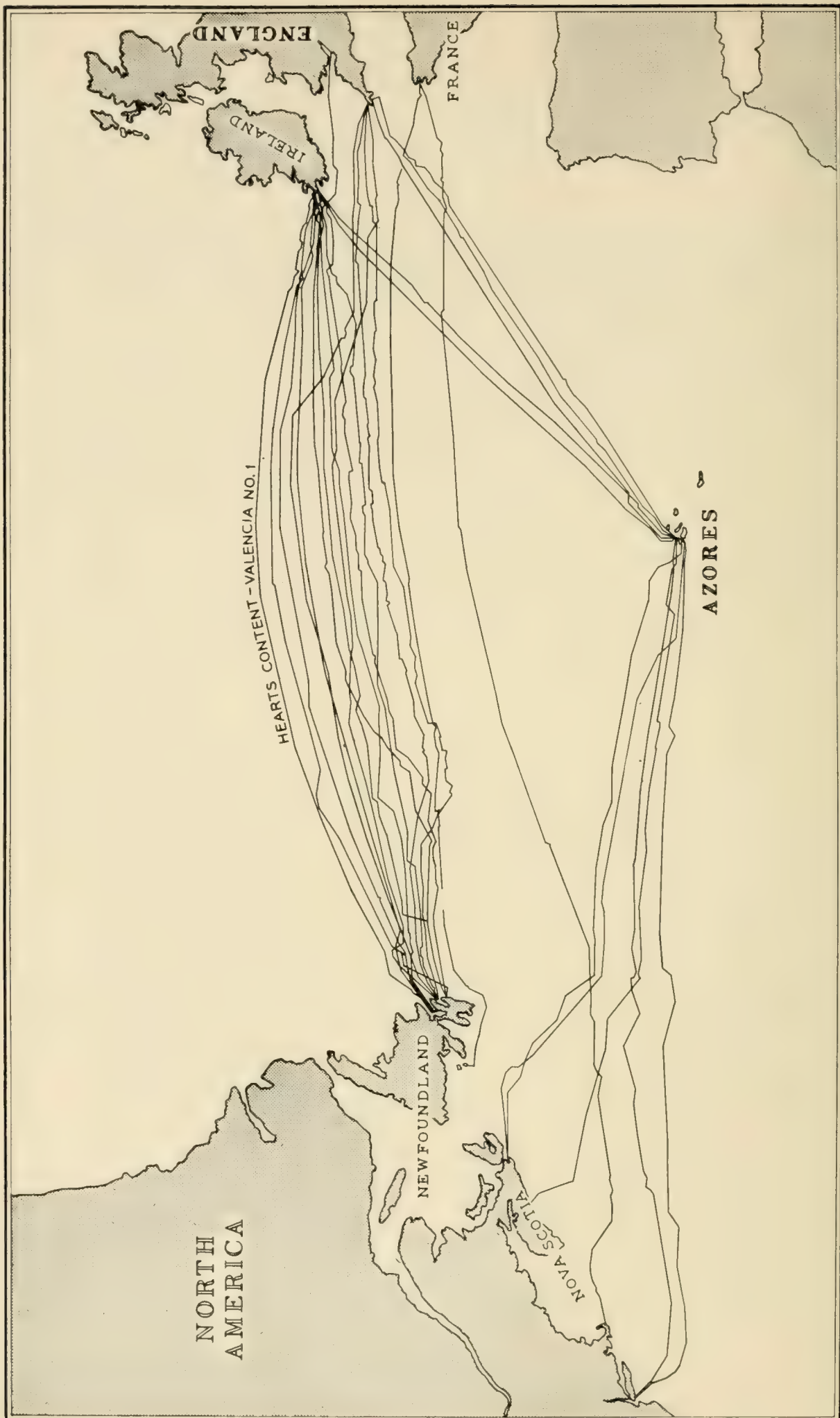


Fig. 1 — Telegraph cables in the North Atlantic.

weather of unpredictable and frequently unpleasant nature. The bridging of this required the laying of two one-way cables over carefully selected routes, using an available cable ship. And the presence in these cables of 102 flexible repeaters posed problems quite unique for such long and deep cables, as also did the need for trimming the system equalization during laying so that transmission over the completed system would fall within the prescribed limits.

From Terrenceville, Newfoundland, to Sydney Mines, Nova Scotia, a single cable, 270 nautical mile path was required through Fortune Bay and across Cabot Strait. While this water is considerably shallower, here again a relatively conventional cable laying problem was complicated by the presence of repeaters which in this section were rigid units, 14 in number.* Trimming of system equalization was also required.

These cables were laid during the spring and summer months of 1955 and 1956. And the preparation for the laying required many months of effort in fields which were for the most part quite foreign to the usual scope of land wire telephone activity. Some appreciation of the problems encountered in this phase of the venture may be gained from the following sections.

NORTH ATLANTIC LINK

Route Selection

The first successful transatlantic telegraph cable was laid across the North Atlantic in 1866. Since that date 15 direct cables have been laid and 5 cables by way of the Azores. The approximate routes of these cables are shown in Fig. 1. It is at once evident that the shortest and possibly the best routes were already occupied so that selection of routes for the two transatlantic telephone cables could be expected to present some difficulty.

Some of the more important considerations which guided the selection were (a) route length, (b) clearance for repairs, (c) trawler and anchor damage possibilities, (d) terminal locations suitable for repeater stations, with staffing in mind as well as facilities for onward routing, due consideration being given to the strategic aspects of the locations.

Route Length

Obviously, the shorter the length of a submarine route, the better. In the present instance, any system length much in excess of about 2,000

* Two additional repeaters are located in the 55 nautical-mile section which is trenched in across Newfoundland from Clarenville to Terrenceville.

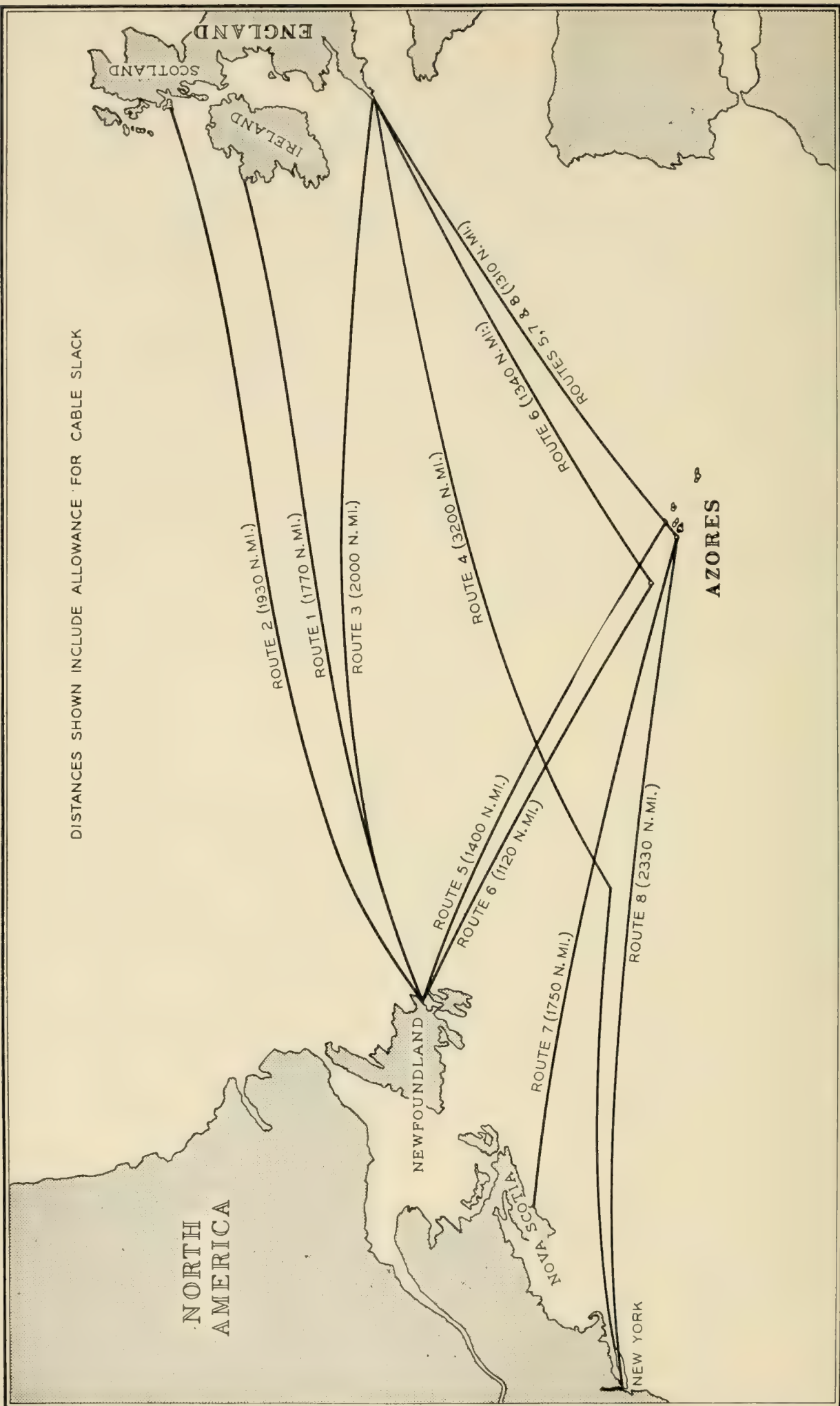


FIG. 2 -- Tentative telephone cable routes.

nautical miles would have resulted in a reduction in the number of voice channels which could be derived from the facility.

On Fig. 2 are shown a number of the routes to which consideration was given in the early planning stages. The distances shown are actual cable lengths and include an allowance for the slack necessary to assure conformance of the cable to the profile of the ocean bottom.

Route 1, from Eire to Newfoundland, at 1,770 miles, is the shortest route and in point of fact was provisionally suggested in 1930 for a new cable. But the difficulty of onward transmission of traffic to London made this route unattractive.

Route 2, from Newfoundland to Scotland, compared favorably in length with Route 1, but its adoption was dependent upon location of a suitable landing site in Scotland.

Route 3, from Newfoundland to Cornwall, England, approximated 2,000 miles laid length and would have been very attractive had not so many existing cables terminated in southern Ireland or the southwest corner of Cornwall, which would lead to a great amount of congestion and consequent hazards to the telephone cables.

Route 4, from New York to Cornwall, was too long to be considered as its length amounted to some 3,200 miles.

Routes 5, 6, 7 and 8 were indirect via the Azores. They were attractive, as only relatively short lengths were involved and suitable sites for intermediate cable stations could have been found on one of the several islands in the Azores. But difficulties attendant upon landing rights, and staffing problems in foreign territory could be foreseen.

Clearance for Repairs

Repair of a faulty cable or repeater necessitates grappling, and in deep water this is likely to be a difficult operation. To avoid imperiling other cables while grappling for the telephone cables and, conversely, to provide assurance against accidental damage to the telephone cables from the grappling operations of others, it was considered essential that the route selected provide adequate clearance from existing cables. Suitable clearance is considered to be 15 to 20 miles in the ocean, with less permissible in the shallower waters of the continental shelves.

Trawler and Anchor Damage Possibilities

It is probable that fishing trawlers cause more interruption of submarine cables than any other outside agency. Cables laid across good fishing grounds are always liable to damage from fouling by the otter

boards of the trawlers. Final splices, either initial or as a result of repair operations, are especially vulnerable to damage because of the difficulty in avoiding slack bights at such points. It was desired, therefore, to avoid fishing grounds if at all possible.

If cables are laid in or near harbors frequented by merchant shipping, damage must be expected from vessels anchoring off shore in depths of less than 30 fathoms and proposed routes should, therefore, avoid such areas.

Cable Terminal Siting

Location of the cable terminal stations must be considered from the standpoints of suitability of shore line for bringing the cables out of the water and also from the standpoint of amenities for the staff. This latter factor is most important in keeping a permanent well-trained staff. For example, owing to staff difficulties, it was necessary to move a terminal station of one company from the west side of Conception Bay in Newfoundland to a site within easy reach of St. Johns.

A further factor in proper siting of the cable terminals is consideration for onward routing of the circuits carried by the cables.

And finally in view of the importance, generally, of submarine cable facilities, it is considered desirable to avoid cable terminal locations in or near a potential military target area and, if at all possible, consideration should be given to underground or protective construction for the terminal stations.

Preliminary Selection

The routes for the telephone cables were considered in the light of the foregoing and after preliminary discussion it was agreed that the two new cables should lie north of all existing cables, should avoid ships' anchorages and should lie on the best bottom which could be picked, clear of all known trawling areas.

In 1930, A.T. & T. Co., in conjunction with the British Post Office, gave serious consideration to the laying of a single coaxial telephone cable between Newfoundland and Frenchport, Ireland (Route 1, Fig. 2). A tentative route was plotted and the cable ship *Dominia* steamed over this taking a series of soundings. These soundings indicated that good bottom was to be found about 20 miles north of the Hearts Content — Valencia cable of 1873. This cable was the most northerly of the telegraph cables spanning the Atlantic. Study of its life history

indicated that faults clear of the continental shelf were few and far between throughout its long life.

The latest British Admiralty charts and bathymetric charts of the U. S. Hydrographic Office for the North Atlantic Ocean were scrutinized and from these and a study of all other relevant data, two routes were plotted which appeared to fulfill the necessary requirements so far as possible. However, it was agreed that if possible the selected routes should be surveyed so that minor adjustments could be made if desirable.

Landing Sites

East End — It was now necessary to find suitable landing sites having regard for the decision that the telephone cables should be routed north of all other existing cables.

On the British side it was necessary to look north of Ireland.

The North Channel, the northern entrance to the Irish Sea, divides northern Ireland from Scotland and had this channel been suitable, the telephone cables might have been run through it to a terminal station on the southwest coast of Scotland in the vicinity of Cairn Ryan. However, the tidal streams through the channel are strong, at least 4 to 5 knots at spring tides; the bottom is rocky and uneven, with overfalls, and any cable laid through it would have a very short life indeed.

It was therefore necessary to search farther north. The west coast of Scotland presents a practically continuous series of deep indentations and bald, rocky cliffs and headlands. The chain of the Hebrides Islands stretches almost uninterruptedly parallel with and at short distances from the coast. It was obviously most desirable to land the cable on the Scottish mainland, and close to rail and road communication if at all possible.

From previous cable maintenance experience it was known that Firth of Lorne which separates the island of Mull from the mainland was a quiet channel, little used by shipping or frequented by trawlers and with tidal streams which were not strong. Earlier passages of Post Office cable ships through the Firth had yielded a series of echo sounding surveys which indicated that except for a distance of about 5 or 6 miles in the vicinity of the Isles of the Sea, the bottom was fairly regular. Several small bays on the mainland side of the Firth just south of Oban appeared from seaward to be very suitable landing sites and this was confirmed by a survey party, which selected a small bay locally named Port Lathaich for the cable landing and site of the station.

The fore shore was mainly firm sand with outcroppings of rock which could be avoided easily when landing the cables. The seaward approach

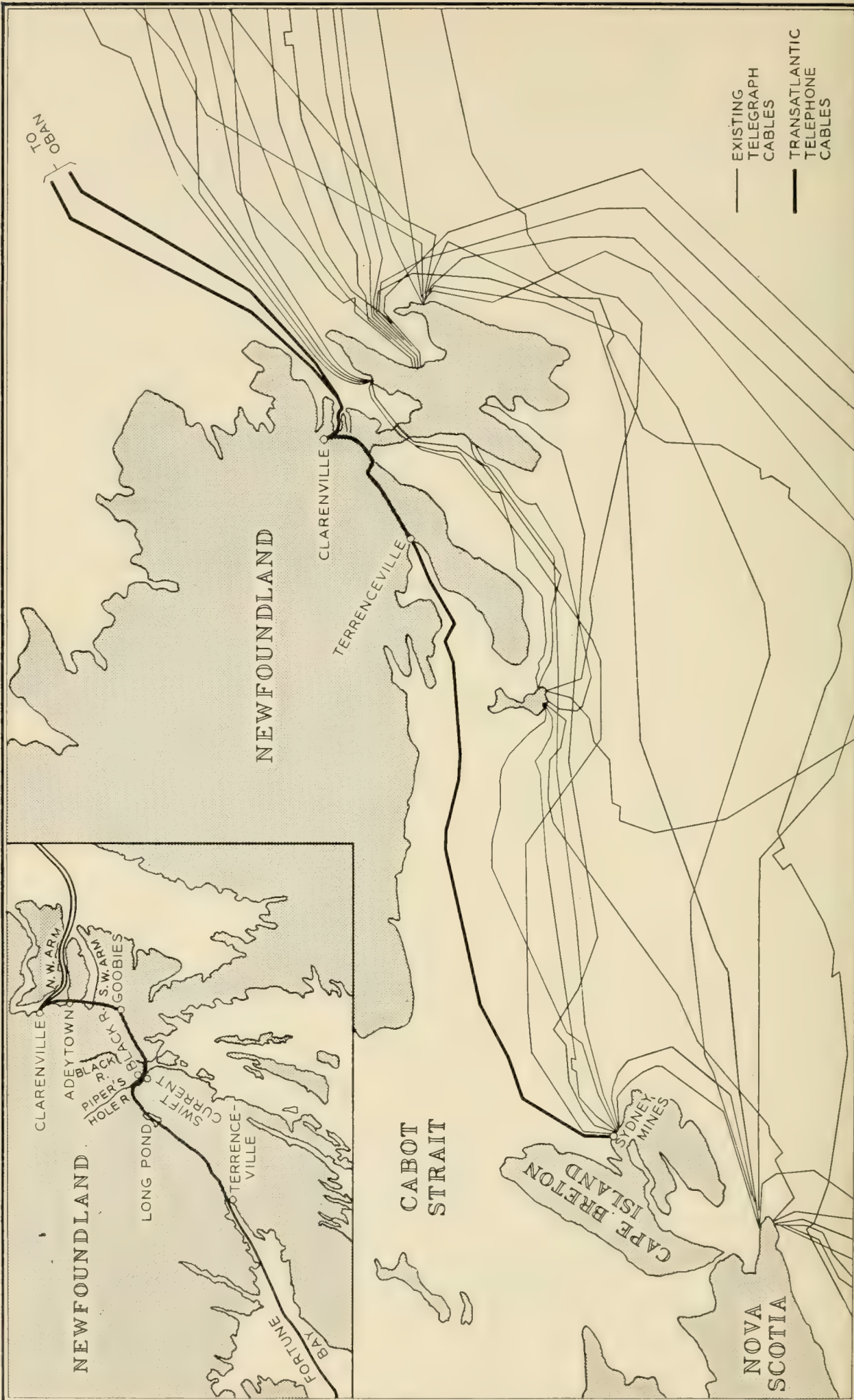


FIG. 3 — Cable landings on Newfoundland, and final route of overland section of the Clarenville-Sydney Mines cable.

was clear of danger and there was ample room to land two cables with a separation on the shore of some 30 yards.

Port Lathaich is only about 3 miles by road from Oban. Additional land cables would be necessary, however, to carry traffic to the main trunk network. From a strategic point of view, although Oban might only just be considered a target area, the cable landing was sufficiently remote to be relatively safe, especially if the cable terminal station was sited in the rocky hillside. To ascertain whether any serious chafing or corrosion would result if cables were laid over the uneven bottom in the Firth, some 8 miles of coaxial cable with E type armoring were laid over the area and recovered after 2 years. There was no evidence of any chafing or corrosion. It was therefore decided that the telephone cables should be routed through the Firth of Lorne to the cable terminal station site at Port Lathaich.

West End — The choice of a suitable cable landing in Newfoundland was more difficult to make, in view of the rugged and sparsely populated nature of the country. From Fig. 3 it will be seen that the existing telegraph cables spanning the Atlantic land either just north of St. Johns, in Conception Bay, or in Trinity Bay. North of Cape Bonavista the coast becomes more broken, and the sea approach is not good. Accordingly, there was no good alternative to routing both telephone cables into Trinity Bay, close to and northwest of the telegraph cable landing at Hearts Content on the southern shore of the bay. A survey party made an extensive examination of all likely places on the western side of the bay from Cape Bonavista in the north to Bull Arm at the southern end where, incidentally, the first successful telegraph cable was landed. Careful consideration of all of the places visited led to the agreement that Clarenville was the best site for a landing and for a cable terminal station.

Clarenville is at the head of the Northwest Arm of Random Sound. It is a junction on the main railway, and a good road to St. Johns will pass through the town in traversing its course from St. Johns to Port aux Basques. Clarenville has a growing population of some 1,600 inhabitants, with stores and repair facilities of various sorts. Good cable landing sites are available just out of town and the approach from the sea up the Northwest Arm presents no navigational difficulties. Such few small vessels as ply to Clarenville during the summer months are not likely to interfere with the cables.

Final Route Agreement

Having agreed Clarenville, Newfoundland, and Oban (Port Lathaich), Scotland, for shore terminations, it was possible to complete the routes

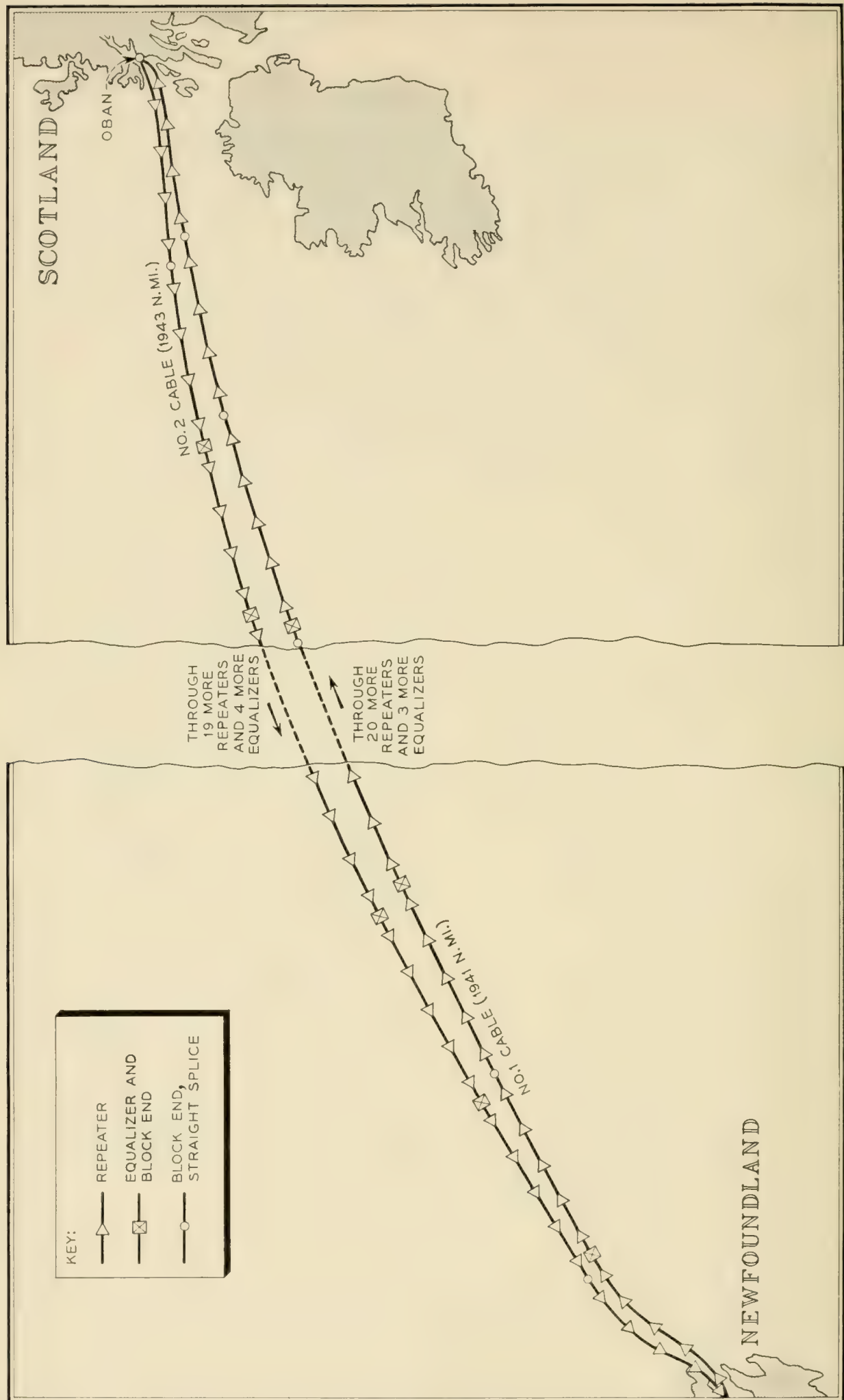


FIG. 4 — Transatlantic telephone cable routes.

for the two cables as shown in Figs. 3 and 4. The final routes are clear of existing cables and avoid crossing known trawling areas and anchorages. The cable stations are well sited with regard to staff amenities, accessibility and strategic requirements. Soundings taken during the laying of the two cables showed a very even bottom except in the Firth of Lorne and one or two places in Trinity Bay. The general profile of the route is shown on Fig. 5.

It is considered that these routes have been selected with care and meet all of the requirements of a well planned cable project. Time alone will tell how well the objectives have been met.

Cable Laying

Early Methods

In 1865 when the legendary *Great Eastern* was pressed into service to lay the first successful transoceanic telegraph cable she was fitted out with certain special cable handling gear. The need for such gear had been amply demonstrated by events which transpired during two earlier and unsuccessful attempts by *H.M.S. Agamemnon* and *U.S.S. Niagara*.

For her assignment, *Great Eastern* was fitted with three large tanks into which her cargo of cable could be coiled. She was also provided with a large drum about which the cable could be wrapped in the course of its passage from the tanks to the sea. This drum was connected to an adjustable braking mechanism which provided the drag necessary to assure that the cable pay-out rate was correct with relation to the speed of the vessel. In addition, a dynamometer was provided so that the stress in the cable would be known at all times. A large sheave fitted to the stern of the ship provided the point of departure of the cable in its journey to the sea bottom.

On Friday, July 13, 1866, *Great Eastern* steamed away from Valencia, Ireland, and 14 days later, on July 27, she arrived off Trinity Bay, Newfoundland, and completed the landing of the western shore end.

H.M.T.S. Monarch

Early in the planning for the transatlantic project it was realized that in no small measure the success of the venture would depend on availability of a vessel suitable for laying the cables. It was fortunate that one of the partners to the enterprise was also the owner and operator of the largest cable ship in all the world, and one well suited to the task at hand.

The twin-screw cable ship *Monarch*, Fig. 6, was built for H.M. Post Master General by Messrs. Swan, Hunter and Wigham Richardson, Ltd., at their Neptune Works, Walker-on-Tyne. She was completed in 1946. This ship is of the shelter deck type having principal dimensions as follows:

Length overall.....	482 feet 9 inches
Breadth moulded.....	55 feet 6 inches
Depth moulded to shelter deck.....	40 feet 0 inches
Gross tonnage.....	8,056

The ship has an overhanging bow which carries three cable sheaves, a cruiser stern with the after paying out cable sheave offset on the port quarter, a semi-balanced rudder having extra large surface, and a cellular double bottom extending from the collision to the aft peak bulkheads. Both main and shelter decks are steel and extend her complete length.

The cable is carried in four welded steel cable tanks fixed to the tank top plating. These are arranged along the ship's center line in a fore and aft direction forward of the main propelling machinery space. They are each 41 feet in diameter and have the following cubic capacities:

	Coiling Space	Gross Cubic Feet
No. 1 Tank.....	33,730	40,170
No. 2 Tank.....	31,820	38,460
No. 3 Tank.....	30,865	37,375
No. 4 Tank.....	30,230	36,300

The opening in the shelter deck above each tank is a circular hatch 8 feet in diameter.

A water tight cone of steel plates is built in the center of each tank to insure against fouling of the cable during payout. Further control

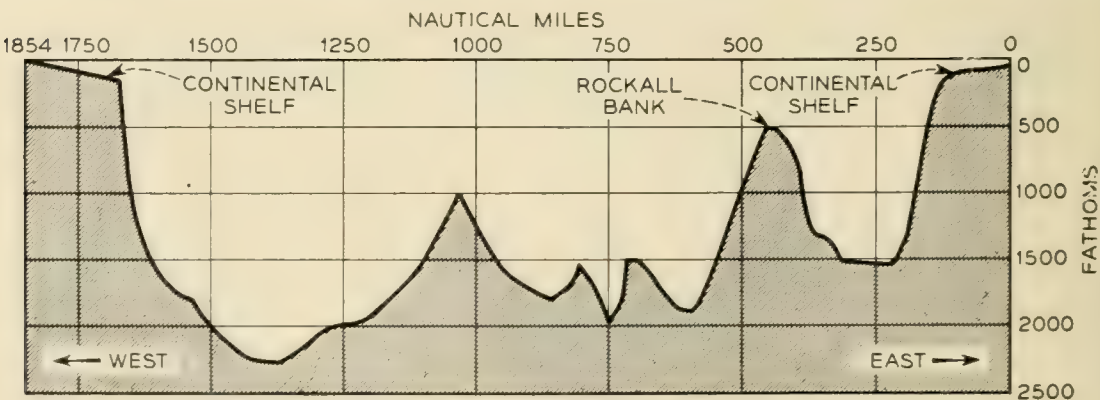


FIG. 5 — Profile of ocean depths between Clarenville and Oban.



FIG. 6 — *H.M.T.S. Monarch.*

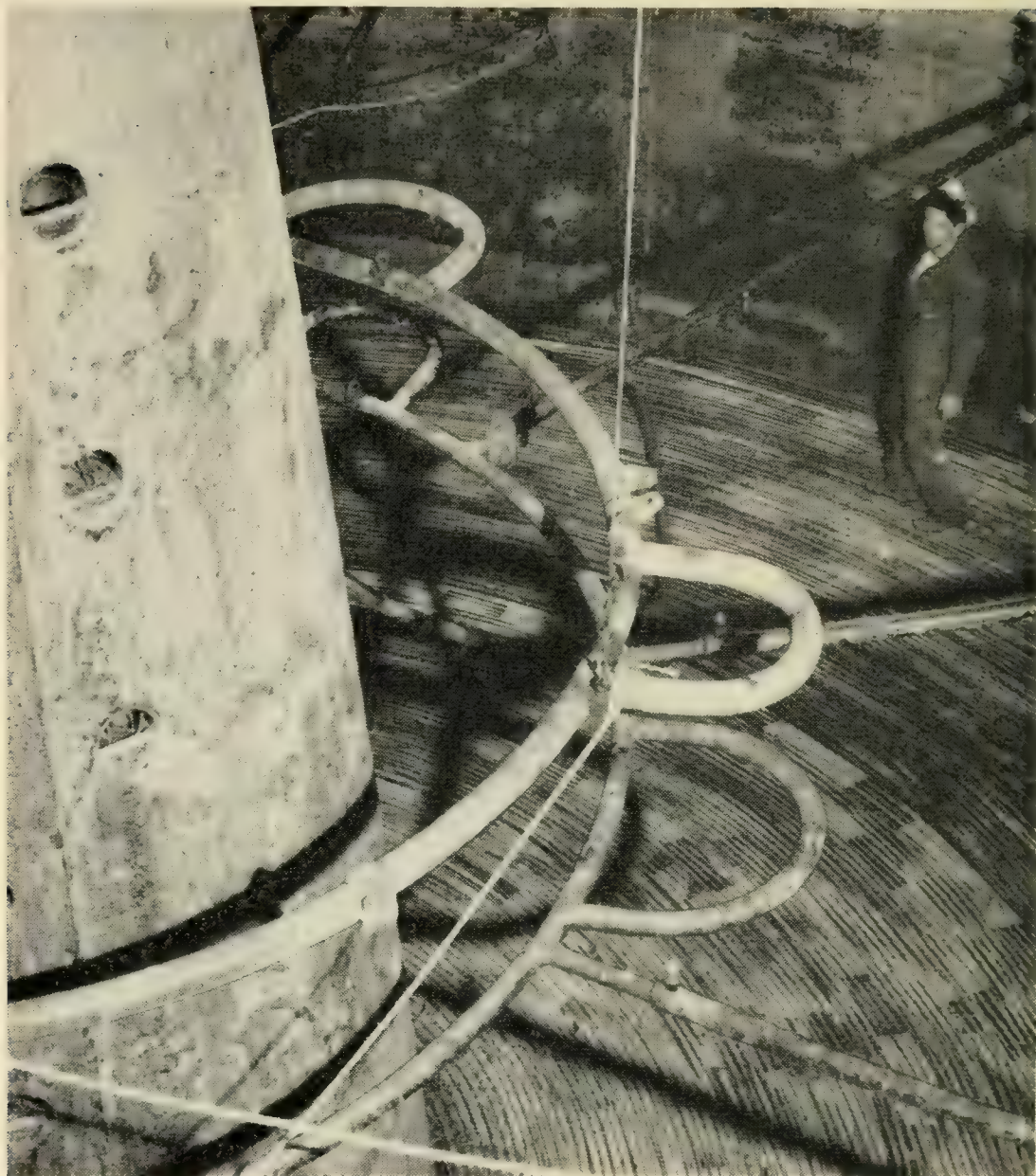


FIG. 7 — Interior of cable tank showing central core, crinoline and flake of cable.

of the cable is provided by a crinoline, Fig. 7, which is a circular spider of steel tubing normally suspended from 1 foot to 3 feet above the top layer of cable in the tank. The crinoline tends to prevent flying bights of cable and also provides a safety platform, in case of trouble, for the men who work in the cable tanks. Each crinoline may be raised and lowered by an electric motor drive.

The maximum cable carrying capacity is approximately 5,000 long tons, or almost 2,000 miles of the deep sea type of cable used on this project if no repeaters were involved.

Monarch is driven by two steam engines. The maximum propeller revolutions are estimated at 110 per minute, giving a ship's speed of about 14 knots.

Two cable engines are fitted forward, both capable of being used for picking up or paying out. These are driven by electric motors having a maximum rating of 160 hp, which will permit picking up at a rate of 0.9 nautical miles per hour with a stress of 20 tons, or at 3.5 knots with a stress of 5.3 tons. The drive system is constant current, so designed that a uniform torque may be held at the drum for any setting of the speed control. When paying out, these motors operate as generators to provide electrical braking, and auxiliary mechanical brakes are also provided.

A single cable engine is fitted aft and this is the main paying out gear. In addition to the electrical brake, the aft engine is provided with a multiple drum externally contracting band brake, manually adjustable and water cooled, and with a further auxiliary fan brake. The fan shaft is driven in such a manner that when cable is being paid out at approximately $8\frac{3}{4}$ knots the fan will revolve at 1,000 rpm and absorb 120 bhp. Adjustments in this are effected by varying the amount of opening in the fan shroud so that as little as 27 bhp may be absorbed.

Dynamometers, both fore and aft, provide for measurement of the cable tension.

Taut wire gear is furnished on the starboard quarter to provide an effective means for calculating the amount of slack paid out. With this gear, steel piano wire, anchored to the bottom, is paid out at constant tension and provides a rough measure of distance steamed over the ground.

A test room with trunks to each cable tank is provided on the shelter deck and fitted with instruments for measuring and locating faults.

Modifications for Flexible Repeaters

In the normal cable-paying-out process, the cable is drawn from the tank, carried along fairleads to the holdback gear (a mechanism for applying slight tension to the cable so that it will snub tightly around the drum), and then wrapped around the drum of the cable engine from two to four turns depending upon the weight of the cable and the depth of the water. At the drum, a fleeting knife is fitted which pushes over the turns already present to make way for the oncoming turn. From the drum the cable passes through the dynamometer and thence to the overboarding sheave.

The Bell System repeaters, manufactured by the Western Electric

Company, were designed with the objective of making them act as much like cable as possible.¹ Despite this, their presence introduced a loading and laying problem as their ability to bend without injury is limited to about $3\frac{1}{2}$ ft radius, and their structure is such that unnecessary bending may involve a hazard to their water tightness. As the majority of the sheaves and drums of the conventional laying gear are considerably smaller than 7 ft diameter, a number of modifications were required in *Monarch's* equipment to satisfy the repeaters.

For the most part, the new gear was designed by the Telegraph Construction and Maintenance Co., Ltd., to broad requirements supplied by the A.T.&T. Co. The modifications included providing the port bow sheave with a flat tread to bring its diameter to 6'10", and replacing both forward dynamometers and the aft dynamometer by a new design employing a 7-foot wheel in a pivoted "A" frame bearing on Elliott pressure type load cells. Port and starboard forward drums were replaced with the maximum diameter drums possible without a major change in the complete gear. This diameter proved to be 6'10" on the tread. The after paying out drum was replaced with one having a 7'0" diameter. The forward port and aft cable drums were equipped with

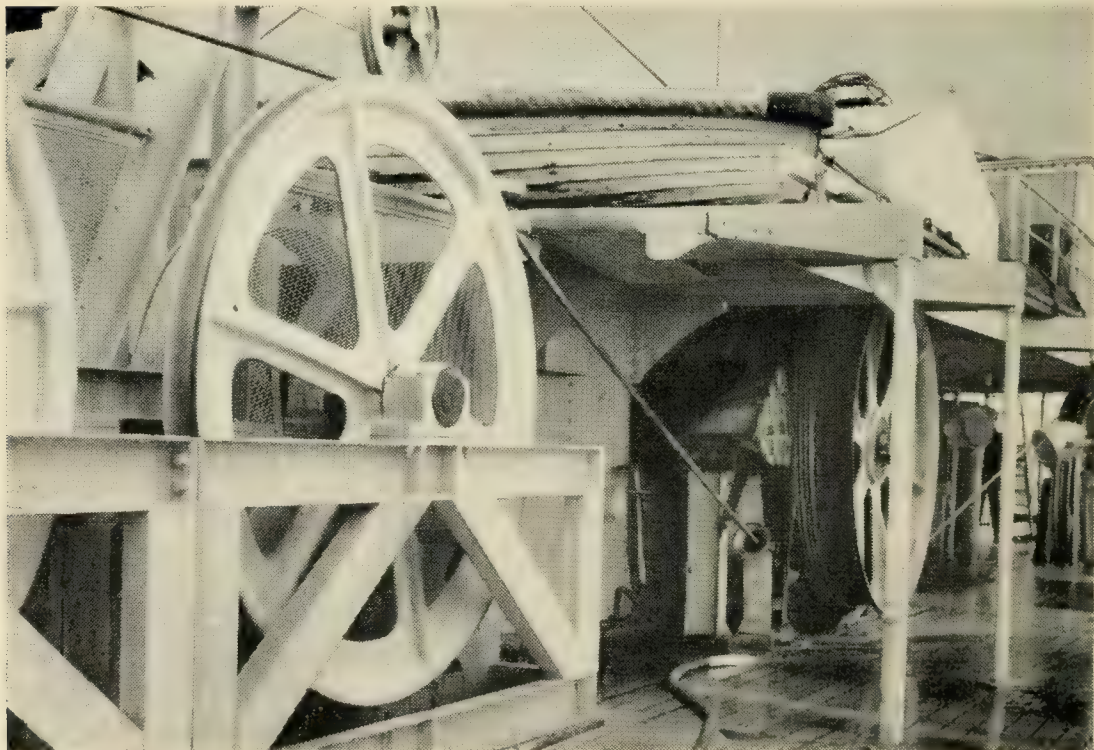


FIG. 8 — General view of modified after cable gear (one of 2 hold back sheaves, drum with fleeting knife and ironing board, and, at extreme right, dynamometer sheave).



Fig. 9 — Cable payout over the stern.

ironing boards. (An ironing board is a curved shoe placed adjacent to the cable drum and spring loaded so that it will force the repeater to conform to the curvature of the drum as it goes on.)

The forward port and starboard draw-off gear sheaves were replaced with larger ones 7'0" in diameter which were made traversable. The aft hold-back gear, of the double sheave type, was also replaced with units having 7'0" sheaves.

Fig. 8 shows a general view of the modified after cable gear, and the 7-ft stern sheave may be seen in Fig. 9. A line schematic of the gear will be found on Fig. 10.

Roller type fairleads shaped into arcs of minimum $3\frac{1}{2}'$ radius were fitted at each cable tank hatch, with smaller roller guides at convenient points to assure fair lead of the cable from the tanks to the cable machinery. Electric hoisting gear was provided for the crinoline in each tank as it was necessary to raise the crinoline whenever a repeater left the tank.

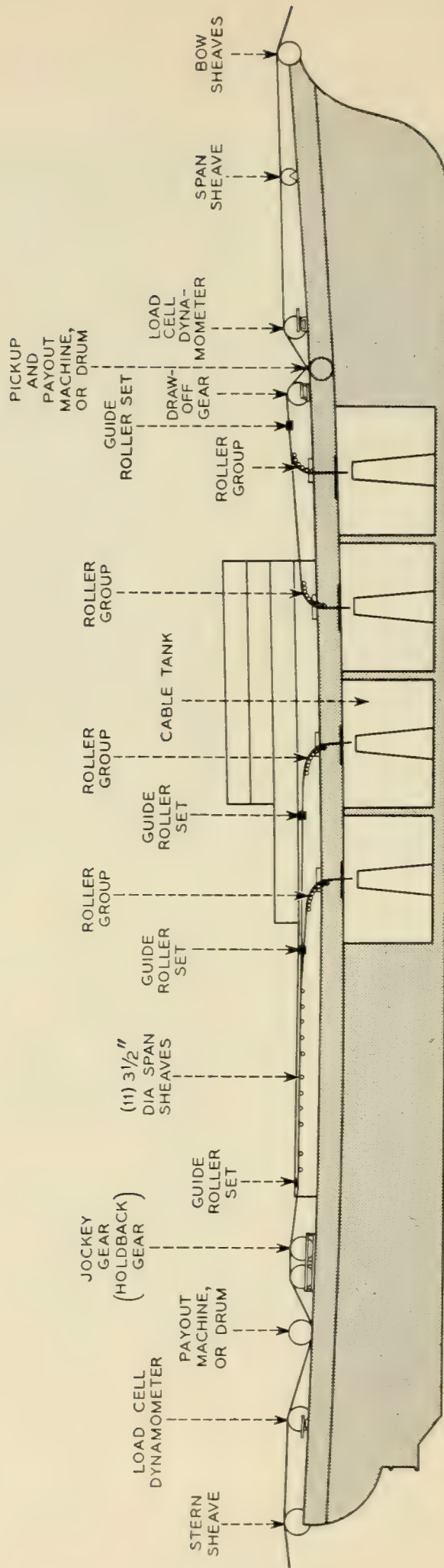


Fig. 10 — Line schematic of the cable gear on *Monarch*.

The test room was greatly enlarged and fitted with the special gear necessary for powering and measuring the system during laying.

Loading Considerations

When the ship is loaded, the cable is coiled carefully in the tanks, layer upon layer — each layer being called a “flake”. The coiling is started from the outside of the tank and progresses clockwise toward the center so that the armor is untwisted one revolution for each complete turn in the tank. When paid out in the reverse order, this twist is restored.

Handling of the repeaters during loading presents a problem because of the need to restrict their bending. After some experimental work was carried out, splints were devised to provide the needed rigidity. These consisted of two angle irons each 12 ft long and equipped at the ends with cold rolled steel rods ranging in length from $1\frac{1}{2}$ ft to 6 ft. By this device it was possible to maintain rigidity over the main central portion of the repeater, including the junction of the core tube with the end nosing, and provide limited flexibility along the outer ends of the core tubes which are less sensitive to bending. The splints were removed once the repeater reached the tank.

Repeaters are always stowed at the outside of the flake where they need be subjected to only a minimum of bending. They are protected with wood dunnage, which must be removed before the repeater is paid out.

With these modifications all repeaters and equalizers were laid successfully from either forward or aft gear at a cable speed of around three knots.

Testing and Equalization

Purpose — Once a submarine cable system has been installed, it is accessible only at the ends for adjustment to improve performance, save at great difficulty and large cost. As some irregularities cannot be corrected from the ends, it behooves the designers to discover and account for such irregularities and to correct them before the cable is finally on the ocean bed.

The laying period offers the last opportunity for accomplishing this, and indeed all too frequently, also the first. This fact, coupled with the broadband design of the link and with the presence in the system of active elements (the repeaters), necessitated a very comprehensive program of tests and measurements during laying.

The program had three specific purposes; (1) to detect immediately any fault which might develop during laying; (2) to permit the design and execution of corrective system adjustments while en route so that transmission performance of the completed link would fall within specified objectives; and (3) to gather data on system characteristics at intermediate points for eventual use in fault location or in aging studies. The need for parts (1) and (3) of the program is more or less self evident, but part (2) merits some further discussion.

In an ideal submarine cable system in an average environment, the attenuation of the cable from one repeater to the next would be offset exactly across the frequency band by the gain of the following repeater. Such a result is never achieved in practice, as the temperature and pressure environments (which affect cable attenuation) cannot be known precisely in advance, and the cable structure itself cannot be manufactured for mile after mile without variation in transmission characteristic. Additionally, the mechanics of the laying process induce minor changes in the physical structure of the cable which reflect in attenuation changes.

If such deviations from desired characteristic produced only differences from the specified system gain objective, compensation could be readily applied at the ends of the submarine link. Unfortunately, this is only partly the case. Their more important effect is the resulting misalignment of operating levels of individual repeaters from design objective.

Misalignment magnitudes must be watched carefully, for at best misalignment narrows the system latitude for seasonal temperature changes and for aging, and at worst it can result in intolerable system noise. If a repeater is preceded by too much cable the signal to noise ratio at the repeater input will be less than desired because of thermal noise. In the opposite case of too little cable, the signal level will be too high and the resulting overloading in the repeater will also affect the signal-noise adversely. Once present on the signal, the noise cannot be removed, and so the cure for excessive misalignment must be applied before the misalignment has developed. Adjustments at intermediate points along the route must therefore be contemplated.

Testing Program — The program which was evolved to meet the three objectives outlined was meticulously reviewed and practiced before the start of laying, and various forms were prepared for entering data and plotting and evaluating results. This was essential to avoid wasting effort or missing valuable data. The wisdom of this was fully apparent to all involved after experience with the close time schedules and the mental tensions which developed during the actual laying.

Staffing for testing was provided by crews of 2 or 3 trained engineers located at the transmitting cable station and on shipboard. Those on the ship served $4\frac{1}{2}$ hour watches at 9 hour intervals, which permitted a reasonable amount of rest and avoided continuous "dog watch" duty by any one crew.

Close contact between shipboard and cable station crews was essential, and was achieved by means of cable and radio order circuits (or "speakers"). Communication from shore to ship when the cable was powered made use of the standard cable order wire circuit at the cable station to apply a signal in the frequency band 16–20 kc. The signal was demodulated and amplified aboard ship by a special stripped version of this same gear. The radio order circuits employed special land antennas and equipment, and for the most part the ship's standard single sideband telephone set, although other equipment at medium frequency was sometimes used for short distances. Radio telegraph with hand keying was available for back-up when conditions were too poor for the radiophone sets.

Plans called for powering the cable at all times except when splices were being made. This was necessary for the measurement program, of course, but also provided additional assurance of safe laying of repeaters, as the glassware and tungsten heaters of the vacuum tubes are more resistant to damage when hot. Power for the first half of each crossing was provided from the cable station. Beyond this point, the required voltage would have become excessive and so the shipboard supply was inserted into the series power loop and its voltage adjusted in proportion to the amount of second half cable actually in the loop.

Monitoring against the possibility of faults was accomplished by measurement of a pilot tone at 160 kc, transmitted over the cable at all times except when data were being taken or power was turned down. Audibly alarmed limits were set on the measurement to indicate any significant deviation in transmission. In actuality, all unanticipated received alarms were found to have resulted from frequency or voltage shifts in the primary shipboard supply for the measuring equipment.

During the design of the system,² consideration of the misalignment problem had indicated the desirability of splitting the cable for each crossing into a number of sections, called ocean blocks. These contained either 4 or 5 repeaters, and were 150 to 200 miles in length. In loading the ship, the two ends of each ocean block were left accessible for connection to the test room and for splicing operations.

Measurements made in the spring of 1955 off Gibraltar had indicated an unexpected change in attenuation called "laying effect",³ which required some last minute adjustment of the repeater section lengths.

With incorporation of these changes, it was known that the factory lengths of cable between repeaters were adequate to keep misalignment within an ocean block within reasonable limits. The system could then be equalized between ocean blocks so that the signal level at the first repeater of a new block would be approximately correct, and the total system noise thus would fall within limits.

This equalization was accomplished in two ways. Excess cable of the order of $\frac{1}{2}$ to 3 miles in length was provided at the top end* of each ocean block. Based on measurements, this could be cut longer or shorter than the nominal spacing of repeaters, so that the repeater gains and cable losses would be matched at some frequency in the band. Residual deviations in other parts of the band could then be mopped up if necessary by inserting a simple equalizer, housed in a container similar to those used for the repeaters. Ten such equalization points were provided in each cable.

In practice, sending levels were adjusted at the cable station to give test tones at the grids of the output tubes of the repeaters which if the system equalization were perfect, would be flat across the frequency band, and at the proper level. These tones were measured on shipboard at the end of the ocean block being paid out. The results were plotted against mileage, with one sheet for each frequency being measured. Because of the "laying effect" and of temperature and pressure changes on the cable as it progressed to the bottom, these plots displayed a slope.

The value of loss (or gain) to be ascertained for each frequency was that which would exist when the entire ocean block was on the bottom. To obtain this, it was necessary to extrapolate the curves to the mileage point representing the end of the block in question. The extrapolation was required to avoid stopping the ship at the end of the block, and so had to anticipate the time needed for turning over and cutting the cable end at the proper point, and making one or two splices (depending on whether or not an equalizer was inserted at the point in question).

Having read the extrapolated values from the curves, these were compared with objectives for that block junction, and the deviations plotted. Transparent overlays, showing the net effect of each of several types of equalizer combined with varying amounts of cable around the nominal spacing, greatly facilitated the final decision as to cutting point and equalizer choice.

This implementation of the system undersea equalization represented a very large part of the effort required of the testing crews during laying.

Additional data gathered for fault location, aging studies and other

* First end out of the cable tank.

purposes included precise determination of repeater crystal frequencies on the bottom, gain frequency runs to show up any fine grained structure which might exist in the band, and values of line current and driving voltages.

Copper resistance and capacitance measurements proved to be of dubious value; in the first case because of the temperature/resistance characteristic of the vacuum tube heaters; in the case of capacitance, probably because of polarization effects in the castor oil capacitors used in the repeaters.

Shipboard Test Equipment — A new test room had been equipped for making the above measurements with transmitting and receiving transmission measuring sets² including the crystal test panels. These sets were provided in duplicate to forestall difficulty should one set develop trouble during the laying. The transmitting consoles were required only for use in calibrating the receiving sets, and for some measurements which were made on individual ocean blocks in the ship's tanks.

Additional gear in the test room included a cable current power supply,⁴ and a "Lookator" which is a pulse echo type of fault locator useful from a point in the cable to the adjacent repeater on either side.

Laying Sequence

H.M.T.S. Monarch is the largest cable ship afloat, with capacity for about 2,000 miles of the Type D deep sea cable in her tanks. However, because of the inherent limitation on their bending radius, the presence of flexible repeaters in the cable puts a restriction on the height to which the coil can be permitted to rise in the tanks. For repeatered Type D cable, therefore, *Monarch's* capacity is cut back to about 1,600 nautical miles.

Types A and B cable, used in shallower waters, are considerably larger and heavier than Type D and consequently, less of these can be carried.

The ideal laying program would have involved one continuous passage across the North Atlantic from cable station to cable station. However, this would have required carrying over 1,900 miles of cable including about 300 miles of Type A and something less than 50 miles of Type B. Such an amount of cable would greatly exceed the ship's capacity.

Each cable was, therefore, laid in 3 sections. The No. 1 cable (southernmost), which transmits from west to east, was laid in the following sequence: Clarendville to just beyond the mouth of Trinity Bay, a distance of 200 miles; thence about 1,250 miles to Rockall Bank (a submerged

plateau); and finally the remaining 500 miles from Rockall Bank to Oban. The No. 2 cable followed the opposite sequence, starting at Oban and proceeding in 3 sections of 500 miles, 1,250 miles, and 200 miles to the terminal at Clarenville. Shore ends, about $\frac{1}{2}$ mile long at Clarenville and 2 miles at Oban, were prepared and put in place in advance of the time when they would be needed. At each intermediate point, the cable was "buoyed off" with a mushroom anchor, connecting lines, and a surface buoy of size appropriate for the water depth.

The mileages indicated are actual cable lengths. They exceed the geographical distances between the points involved because of the slack allowance which experience has shown to be necessary to assure reasonable conformance of the cable with the contour of the ocean bed. Normally, about 5 per cent slack is considered desirable in deep water, with the allowance decreasing in steps to zero in shallow water.

All available information indicated that the most favorable weather conditions in the North Atlantic could be anticipated in the period May through August. Prior to May, ice could be expected along the western sections of the route and after August, hurricanes were likely, and later the winter storms.

The laying of the No. 1 cable was started June 28, 1955, and completed September 26. The actual laying period took in but 24 days in this interval, the remainder of the time being spent in transit and in re-loading. The No. 2 cable was started June 4, 1956, and completed August 14. About 16 of laying days were involved.

The routine aboard ship during laying consisted in passing out cable at the rate of 6 to 7 knots for a repeater section length of a little over 37 nautical miles, then slowing down to about 3 knots as the repeater passed through the cable machinery and overboard, then back to speed again. During all of this period the testing crews, both on shipboard and at the transmitting cable station, were busy measuring, recording data and planning the equalization trimming. At a point shortly after the passage of the next-to-last repeater in an ocean block, special measures were required for the equalization program. From that point until the joints associated with the connection to the following ocean block had been completed, the speed was reduced to 5 knots. The need for this arose from the following considerations.

Stopping of the ship in deep water introduces serious possibility of formation of kinks in the cable, and is to be avoided at all costs. To permit continuous laying, it was necessary to determine the amount of cable needed for equalization, measure out this cable, and complete the splices before reaching the end of the block being laid.

The addition of an equalizer at the end of the block requires two joints and armor splices. Preparing the cable ends, brazing together the center conductor and associated tapes, injection molding the polyethylene around the center conductor, replacing and overlaying the armor wires and binding the splice consumes 6 to 7 hours for a single splice, and 8 to 9 hours for two splices when overlapping of operations is practical. An allowance of 3 hours is considered necessary for remolding in event of a defective joint (disclosed by X-ray inspection). The time allowance required to complete the splicing of ocean blocks is therefore 9 to 12 hours. About $1\frac{1}{2}$ hours are needed to carry out the extrapolation, make the equalization decision and turn over cable to the cutting point. During the interval between the cable cut and the completion of the joint, the ship's speed was maintained at 5 knots so as to minimize the distance over which extrapolation of equalization data had to be extended. Even so, the final extrapolation covered the last 60 to 75 miles of each ocean block.

During the jointing intervals, system power was turned down to avoid any hazard to the members of the jointing crew. It was restored as soon as the moldings had been X-rayed and the outer or return tapes had been brazed. These activities were so timed that in almost every case the system was powered as each repeater went overboard.

CLARENVILLE-SYDNEY MINES LINK

Route Selection

Clareville having been selected as the site of the cable terminal station on the west end of the ocean crossing, it was necessary to consider how the system was to be extended to Nova Scotia for connection with the North American continental network.

A number of alternatives were possible as described below:

Alternative 1 contemplated radio relay across Newfoundland to Port aux Basques, and thence across the Cabot Strait. However, a survey revealed that maintenance access to suitable sites would be most difficult, particularly in winter, and primary power was not obtainable.

Alternative 2 involved a poor submarine route around the Avalon Peninsula to possibly Halifax, Nova Scotia. The length of the sea cable would be about 600 nautical miles. It would be necessary to cross many working telegraph cables, (Fig. 3). Trawler damage could be expected as the cable would need to traverse known trawling areas, and during the winter months any repairs would be costly and prolonged. Also it was

not desired to lay another cable out of Trinity Bay as the route might be wanted for a second transoceanic cable at some time in the future.

Alternative 3 involved a submarine cable from Clarenville out through the North West Arm to Rantem at the head of Trinity Bay, a short land cable across the isthmus, and thence either a submarine cable direct to Sydney, Nova Scotia, or a land crossing of the Burin Peninsula at Garnish and thence by submarine cable to Sydney. This route involved three open sea sections with one or two land sections. There was rather limited space for a cable in Placentia Bay and the bottom was uneven and rocky. Existing cables laid around the Burin Peninsula have had interruptions which indicated an unsuitable bottom and fishing trawlers had been seen in the vicinity recently.

Alternative 4 also involved a cable overland, from Clarenville to Terrenceville at the head of Fortune Bay, there to join a direct submarine cable to Sydney Mines. Three short underwater sections would be involved in the Clarenville-Terrenceville link, but these could be in shallow water out of harm's way. The main submarine route from Terrenceville to Sydney Mines would be clear of other cables and would avoid trawling areas and anchorages, a not inconsiderable achievement in view of the congestion of submarine cables and the fishing activity around the southeast corner of Newfoundland. Further, a good landing site in Nova Scotia was available on property near Sydney Mines owned by Eastern Telephone and Telegraph Company.

After due consideration, *Alternative 4* was chosen as being the most satisfactory from all aspects and the final route is shown on Figure 3.

This route is considered to be most likely to have a good life history. While it would have been possible to have one continuous land cable between Clarenville and Terrenceville, the three short underwater sections saved a considerable amount of trenching without adding undue hazard to the system.

Clarenville-Terrenceville Route

It having been decided to route the Post Office single cable system overland from Clarenville to Terrenceville a number of other matters required decision. The first was the type of cable to be employed for this section. Several alternatives were considered, bearing in mind factors such as the type of terrain, access, availability of primary power, possible future expansion of capacity and, of course, interference from static and radio frequency pick-up. The advantages of using standard solid dielectric coaxial ocean cable with submarine-type repeaters were judged

to outweigh all other considerations and left only one problem, namely, shielding from interference.

Up to this time shore ends of submarine cables used by the Post Office were shielded for about a quarter of a mile from shore by a lead sheath insulated from the return tapes of the coaxial by a polyethylene barrier. Experience indicated that such shielding might not be adequate over a long distance on land. The question was resolved by the addition of iron shielding tapes and a plastic jacket to standard submarine cable. The structure is described elsewhere.⁵

Through the use of this robust, wire armored cable and two steel housed submarine repeaters, no limitations from the noise pickup angle were placed on the detailed route selection for the overland section. The first thought was to try to proceed directly across country from Clarenville to Pipers Hole River, saving at least 10 miles over a route which followed the road, and on which advantage might be taken of quite long water stretches into which the cable could be dropped. Black River Pond, for instance, is $4\frac{1}{2}$ miles long. This proposal was abandoned after surveys, because of the very rocky nature of the country and difficulty of access both for construction and any subsequent maintenance, and it was decided to follow the general course of the roads.

It was possible to avoid trenching in the rocky, precipitous cliff country from Clarenville to Adeytown by laying about 6 miles of cable in the water of Northwest Arm. Similar considerations dictated the choice of two miles of cable in the sea across Southwest Arm. Thence the route followed the road, at a distance ranging from 250 yards to more than a mile, as far as Placentia Bay, taking advantage of the larger ponds where possible to avoid trenching.

Reaching the 800 foot high ground beyond the Pipers Hole River from the north of Black River proved quite difficult. Here the road is carved out of the foot of the cliffs as far as Swift Current and the country behind is solid rock. Plans exist for a hydro-electric project involving dams in Pipers Hole River just north of the road crossing and it is naturally not desirable to bury a cable in such a locality. The river estuary itself passing by Swift Current presents only a narrow 6 foot deep navigable channel at low water but it was decided that this could be used for some 6 miles by employing a barge and a shallow draft tug for the laying.

A suitable route out of the basin up through wooded gorges to the top took about a week of very hard going to locate. Thereafter all was plain sailing taking advantage of ponds such as Long Pond (4 miles) and Sock Pond (3 miles) until the route arrived within 6 miles of Terrenceville.

Here it was reluctantly decided to bury the cable in a deep trench on the inner side of the road, as the other side falls sharply to the sea. The desire to avoid roads was due to their instability and the methods used for construction and repair. This road is dirt only, with no foundations, and in this particular section has been known to slide away into the river bed below.

The final length of cable laid was just short of 55 nautical miles.

Cabot Strait Laying

Coaxial submarine cables in which rigid repeaters are inserted cannot be laid with the existing cable laying machinery except by stopping the ship, removing the turns of cable from the drum and then passing the repeater by the drum and restoring the turns. Special equipment is also needed for launching the repeater over the bow sheaves.

Ship Modifications

The following equipment was installed on *Monarch* for the laying of cables carrying rigid repeaters.

A gantry over the bow baulks, consisting of a 22 foot steel beam projecting 6 feet beyond the sheaves, was installed for handling repeaters at the bow. This gantry was fitted with an electrically operated traveling hoist for lifting the repeaters over the bow sheaves and lowering them into the sea. A standby hand operated lifting block and traveller were provided to guard against failure of the power point.

A rubber-tired steerable steel dolly (or trolley) was developed from the chassis of a small car to transport the repeaters from the cable tank hatches to the bow sheaves. These repeaters weigh about 1200 pounds.

Storage racks were built up from steel sections provided with shaped, rubber lined wood blocks and were fitted at each cable tank hatch on the shelter deck. Each rack held 4 repeaters in double tiers. A hand operated lift was furnished for moving the repeaters from the storage racks to the dolly.

A special quick release grip was furnished for use when lifting the repeaters by the electric hoist on the bow gantry. Deflection plates were also fitted on the fore deck around dynamometers and hatches to avoid their fouling the dolly.

Launching Rigid Repeaters

The rigid repeaters were stowed in their racks in the order of their laying. The bights of the cable attached to the ends of the repeaters

were brought up the sides of the cable tanks, secured along the arms of the crinoline and up the sides of the hatch coamings to the deck, clear of the running length of cable and from where the repeater could be drawn forward along the deck on the dolly.

When the time came for a repeater to be laid, speed was reduced and the ship finally stopped head to wind. A 6x3 compound rope from the starboard cable drum was secured to the cable just abaft the bow baulks. Sufficient cable was then paid out until the tension was taken up by this rope. The turns of the running cable were then removed from the port cable drum and the resulting slack cable worked overboard by paying out the starboard drum rope which was holding the tension. When the excess cable had been cleared from the deck, the repeater on its dolly was carefully hauled along the fore deck to the traveling hoist of the overhead gantry. Cable was then drawn from the tank so that the turns could be re-formed on deck and replaced on the port cable drum. The repeater was lifted from the trolley and traversed outboard as soon as it was high enough to clear the bow baulks. It was then lowered to the water's edge and when the tension had again been taken by the cable, the quick release grip was slipped and the starboard drum rope cut. Paying out was then resumed.

Laying Program

On February 1, 1956, *Monarch*, having returned to Ocean Works, Erith, after refitting, commenced loading the various sections of cable to be used for the Terrenceville-Sydney Mines route. The sections were all carefully tested and measured in the Works before loading. The cable ends were clearly marked and dogged together by a length of rope which was not removed until the repeater had been jointed into its connecting sections of cable.

Loading of the cable and splicing in of repeaters was finished by April 10 and the system tested and checked. *Monarch* sailed for Sydney Mines on April 18 and arrived there April 30. The cable station is situated about $1\frac{1}{2}$ miles inland from the shore, with a small lake intervening. A length of Type B, insulated outer conductor, lead covered cable had previously been laid from the station across the lake to a narrow strip of land which separates it from the sea. The joint to the main cable was to be made on this strip. Two medium sized shore based motor boats were used to tow the end of the double armored section of cable from *Monarch* to the shore. During this journey the cable was supported by empty oil drums at close intervals. When the motor boats had reached

shoal water the end of the cable was secured to a landing line and two tractors took over the hauling.

When enough cable was on shore to make the joint and the splice, the barrels were cut away and *Monarch* weighed anchor and paid out this section of double armored cable on the agreed route and buoyed off the end. She then steamed over the proposed track to Terrenceville, taking soundings and sea bottom temperatures as required, and anchored off Terrenceville on May 3.

Preparations for landing the end were at once put in hand and the ship's motor launches towed the end of the cable towards the cable landing, the cable again being supported by empty oil drums. This end was jointed and spliced to a piece of cable which had been laid previously from the Terrenceville cable hut to a sand spit which juts across the head of Fortune Bay, about a mile away.

Upon completion of the splice, overall tests were made from the ship to the Terrenceville cable hut, and all being well, paying out toward the buoyed end off Sydney Mines was begun on May 4.

The first repeater went over about two hours after the start of laying and the others followed at approximately $4\frac{3}{4}$ hour intervals.

On May 7 the cable buoy on the Sydney Mines end was recovered and the end hove inward. After tests in both directions, the final joint and splice were made. This operation was completed on May 9, and on receipt of a signal that all was well, *Monarch* proceeded into harbor at Sydney to land testing equipment, a spare equalizer and other equipment.

Equalization and Testing

The cable had been loaded into the ship in repeater section lengths, so cut that when laid at estimated mean annual sea temperature, the expected attenuation would be 60.0 db at 552 kc. A correction for the change in attenuation of the cable when coiled in the factory tanks and when laid in about 100 fathoms had been determined from tests on two 10-mile lengths of cable, laid off the Island of Skye. The correction amounted to a decrease in attenuation when laid of 1.42 per cent. This was essentially an empirical result, and as the mechanism of the change was not fully understood, a possible further inaccuracy of equalization might arise.

Sea bottom temperatures along the route were obtained from information supplied by the Fisheries Research Board of Canada, but unfortunately, this information was rather meager and varied considerably with locality.

Since the cable equalization built into the repeater differed appreciably from the final determination on laid cable, it was found necessary at a comparatively late stage to introduce an undersea equalizer into the center of the sea section. This was intended to eliminate a flat peak of loss of 3.5 db, expected at about 100 kc. So that the last repeater should not lie too near the beach at Sydney Mines, a network simulating 9 miles of cable was also inserted in the undersea equalizer.

The repeaters were spliced into the cable lengths on board *Monarch* and tests were made at every stage of the buildup of the system. The equalizer was permanently jointed to the first half section of 7 repeaters and left with an excess length of tail which could be cut as desired during the laying operation to further improve the equalization. The first and second halves of the system were temporarily connected through power separation filters so that the whole system could be energized just prior to laying.

The test routine carried out included attenuation measurement at 5 frequencies in each direction of transmission, noise, pulse and loop-gain, supervisory measurements, dc and insulation resistance and capacitance. *Monarch* test room contained, therefore, two sets of terminal equipments similar to those installed at Clarendville and at Sydney Mines.

It was decided to energize the system continuously during the laying except for the few hours when power had to be removed to make the equalizer splice. This enabled a continuous order (speaker) circuit to be operated over the cable and minimized the number of energizing and warm-up periods. The only disadvantage, considered to be slight, was the necessary omission of insulation resistance and capacitance measurement during laying, except in the course of the equalizer splicing operation.

The plan was to lay from Terrenceville in the direction of the high-frequency band and to test the system to *Monarch* during the laying from this shore station. The overland section between Clarendville and Terrenceville, which contained two repeaters, was connected on with appropriate equalization after the submarine section had been satisfactorily completed and tested.

At Terrenceville, after the cable end had been taken ashore and the beach joint completed, the system was energized from *Monarch* with a dc power ground at Terrenceville for the necessary 4 hour minimum warming up period. The first set of routine measurements of the laying operation was then carried out. Thereafter, a complete set of measurements on the Terrenceville half of the system was made after every 10 miles of cable laid. An occasional check set of measurements was also made on the Sydney Mines half of the system in the tanks.

The primary object of these tests was to determine what length of cable should be inserted between the equalizer and the 8th undersea repeater to obtain the optimum system. In practice this resulted in arranging for a length of cable such that the output level of the 8th repeater at 522 kc should be equal to that at the output of the first repeater at the assumed mean annual temperature of 35.1°F.

The estimated length of cable required for this purpose was plotted after each measurement. It became evident soon after laying the 5th repeater (94.6 nautical miles of cable laid) that the linear relation obtained could be extrapolated with adequate accuracy to safely specify a length of 6.06 nautical miles of cable between the equalizer and the adjacent repeater.

This decision on length was taken, and after removing power the equalizer was accordingly jointed in, the operation being completed before it was necessary to pay out the splice. During this period capacitance and conductor and insulation resistances were checked on each half. During the laying of the second half, measurements were made as for the first half. On arrival at the buoyed shore end a final complete set of measurements was made and these suitably corrected for the shore end length, transmitted to Sydney Mines so that the first measurements from Sydney Mines to Clarenville could be checked with those obtained on the ship.

SIDELIGHTS

Weather

Weather is the big question mark in cable laying and repair activities. With few exceptions the transatlantic project was blessed with remarkably good weather. The exceptions were, however, noteworthy.

Heavy snow squalls were encountered off Terrenceville during the operation in that vicinity. At Rockall Bank on the first lay one heavy storm came up as the last repeater in the 1,200-mile section was going over, and made this launching and the subsequent buoying of the end very difficult operations. Both were accomplished successfully as a result of the superb seamanship of *Monarch's* commander, Captain J. P. F. Betson, and his officers and crew.

A second, and worse storm was encountered upon the return to Rockall. This was a manifestation of hurricane Ione, with wind velocities above 100 mph and very high seas. The ship had given up searching for the buoy (later reported drifting more than 500 miles away off the Faeroe Islands) and was grappling for the cable, when the storm hit. Fortun-

ately, the cable had not yet been found, so the ship could head into the wind and ride it out. She was driven many miles off course in the process, and the seas will be long and vividly remembered by all present. Incidentally, the cable was picked up shortly after the storm had moderated.

Generally speaking, the effect of the weather on the engineering supernumeraries on board was not severe, although *Monarch's* stock of dramamine was somewhat depleted by the end of the project.

Miscellaneous Events

At the start of the first transatlantic lay, several icebergs were encountered. One, a small one at the mouth of Random Sound, lay in the planned path of the cable and caused an involuntary, though minor, revision of the route. The others, beyond the mouth of Trinity Bay, were larger but also farther away.

Whales and grampuses got to be common sights, although much film was expended at first by the uninitiated.

An occasional bird rested on the ship far from land, obviously exhausted from its long and presumably unintended journey.

Progress Bulletins

Daily progress bulletins were radioed to headquarters of all partners during the laying.

In addition, because a telephone cable system differs considerably from submarine telegraph cables, the officers and crew were briefed by the engineering personnel as to the repeater structure, the need for equalization and the general objectives of the venture. This proved to be a very profitable move indeed, for the cooperation of all hands was everything that could be wished. As a follow up, daily performance bulletins were posted in strategic parts of the ship so that everyone no matter what his duties, could know just how the evolving system was performing with respect to objectives.

Cable Order Circuit

One way conversation from shore to ship over the cable was possible all the time the repeaters were energized. This was a source of very great satisfaction to the shipboard test crew, as it was concrete evidence that the cable was working, and working well. When power was turned down, the recourse to radio telephone provided a comparison which generally left no doubt as to the future value of the cable.

ACKNOWLEDGEMENTS

When the final splice was slipped into the water of Clarendville Harbor, on August 14, 1956, there was completed a venture quite unique in the annals of submarine cable laying. And in the laying perhaps more than in any other phase of the transatlantic project did the successful conclusion provide evidence of the friendly and harmonious relationships between the different organizations and nationalities involved, and of the close coordination of their efforts.

REFERENCES

1. T. F. Gleichmann, A. H. Lince, M. C. Wooley and F. J. Braga, Repeater Design for the North Atlantic Link. See page 69 of this issue.
2. H. A. Lewis, R. S. Tucker, G. H. Lovell and J. M. Fraser, System Design for the North Atlantic Link. See page 29 of this issue.
3. A. W. Lebert, H. B. Fischer and M. C. Biskeborn, Cable Design and Manufacture for the Transatlantic Submarine Cable System. See page 189 of this issue.
4. G. W. Meszaros and H. H. Spencer, Power Feed Equipment for the North Atlantic Link. See page 139 of this issue.
5. R. J. Halsey and J. F. Bampton, System Design for the Newfoundland-Nova Scotia Link. See page 217 of this issue.

Bell System Technical Papers Not Published in This Journal

BALDWIN, M. W., JR.,¹ and NIELSEN, G., JR.¹

Subjective Sharpness of Simulated Color Television Pictures, J. Am. Optical Soc., **46**, pp. 681-685, Sept. 1956.

BARBIERI, F.,¹ and DURAND, J.¹

Method for Cutting Cylindrical Crystals, Rev. Sci. Instr., Shop Note, **27**, pp. 871-872, Oct., 1956.

BEMSKI, G.¹

Quenched-In Recombination Centers in Silicon, Phys. Rev., **103**, pp. 567-569, Aug. 1, 1956.

BOMMEL, H. E., see Mason, W. P.

BRADY, G. W., see Mays, J. M.

BROWN, W. L., see Montgomery, H. C.

CUTLER, C. C.¹

Instability in Hollow and Strip Beams, J. Appl. Phys., Letter to the Editor, **27**, pp. 1028-1029, Sept., 1956.

DACEY, G. C., see Thomas, D. E.

DAIL, H. W., JR., see Galt, J. K.

DEHN, J. W.,¹ and HERSEY, R. E.¹

Recent New Features for the No. 5 Crossbar Switching System, Commun. and Electronics, **26**, pp. 457-466, Sept., 1956.

DURAND, J., see Barbieri, F.

FARRAR, H. K., see Maxwell, J. L.

¹ Bell Telephone Laboratories, Inc.

FAY, C. E.¹

Ferrite-Tuned Resonant Cavities, Proc. I.R.E., **44**, pp. 1446-1449, Oct., 1956.

FEHER, G.¹

Observation of Nuclear Magnetic Resonances Via the Electron Spin Resonance Line, Phys. Rev., Letter to the Editor, **103**, pp. 834-835, Aug. 1, 1956.

FERRELL, E. B.¹

A Terminal For Data Transmission Over Telephone Circuits, Proc. Western Joint Computer Conf., pp. 31-33, Feb. 7-9, 1956.

GALT, J. K.,¹ YAGER, W. A.,¹ and DAIL, H. W., JR.¹

Cyclotron Resonance Effects in Graphite, Phys. Rev., Letter to the Editor, **103**, pp. 1586-1587, Sept. 1, 1956.

GARRETT, C. G. B.¹

The Physics of Semiconductor Surfaces, Nature, **178**, p. 396, Aug. 25, 1956.

GOLDEY, J. M., see Moll, J. L.

GOLDSTEIN, H. L.,¹ and Lowell, R. J.¹

Magnetic Amplifier Controlled Regulated Rectifiers, Proc. Special Tech. Conf. on Magnetic Amplifiers, pp. 145-147, July, 1956.

GULDNER, W. G.¹

Application of Vacuum Techniques to Analytical Chemistry, Vacuum Symp. Trans., pp. 1-6, 1955.

HARROWER, G. A.¹

Energy Spectra of Secondary Electrons from Mo and W for Low Primary Energies, Phys. Rev., **104**, pp. 52-56, Oct., 1956.

HERSEY, R. E., see Dehn, J. W.

HITTINGER, W. C., see Warner, R. M., Jr.

HOLONYAK, N., see Moll, J. L.

¹ Bell Telephone Laboratories, Inc.

HOSFORD, J. A.³

The Development of Automatic Manufacturing Facilities for Reed Switches, Commun. and Electronics, **26**, pp. 496-500, Sept., 1956.

HOVGAARD, O. M.¹

Capability of Sealed Contact Relays, Commun. and Electronics, **26**, pp. 466-468, Sept., 1956.

IRELAND, G.⁵

Management Development in the Communications Field, Elec. Engg., **75**, pp. 881-884, Oct. 1956

JACCARINO, V., see Shulman, R. G.

JAYCOX, E. K.,¹ and PRESCOTT, B. E.¹

Spectrochemical Analysis of Thermionic Cathode Nickel Alloys by Graphite-to-Metal Arcing Technique, Anal. Chem., **28**, pp. 1544-1547, Oct., 1956.

JONES, H. L.⁶

A Blend of Operations Research and Quality Control in Balancing Loads on Telephone Equipment, Proc. Am. Soc. for Quality Control, June, 1956.

KISLIUK, P.¹

The Dipole Moment of NF₃, J. Chem. Phys., Letter to the Editor, **25**, p. 779, Oct., 1956.

KNAPP, H. M.¹

Design Features of Bell System Wire Spring Relays, Commun. and Electronics, **26**, pp. 482-486, Sept., 1956.

KUCK, R. G.⁵

Microwave Facilities with Built-In Reliability, Commun. and Electronics, **26**, pp. 438-441, Sept., 1956.

KUNZLER, J. E.¹

Liquid Nitrogen Vacuum Trap Containing a Constant Cold Zone. Rev. Sci. Instr., Shop Note, **27**, p. 879, Oct., 1956.

¹ Bell Telephone Laboratories, Inc.

³ Western Electric Company.

⁵ Pacific Telephone and Telegraph Company, San Francisco, Calif.

⁶ Illinois Bell Telephone Company, Chicago, Ill.

LEWIS, H. W.¹

Ballistocardiographic Instrumentation, *Rev. Sci. Instr.*, **27**, pp. 835-837, Oct., 1956.

LOWELL, R. J., see Goldstein, H. L.

LUKE, C. L.¹

Determination of Traces of Gallium and Indium in Germanium and Germanium Dioxide, *Anal. Chem.*, **28**, pp. 1340-1342, Aug., 1956.

LUKE, C. L.¹

Rapid Photometric Determination of Magnetism in Electronic Nickel, *Anal. Chem.*, **28**, pp. 1443-1445, Sept., 1956.

MASON, W. P.,¹ and BOMMEL, H. E.¹

Ultrasonic Attenuation at Low Temperatures for Metals in the Normal and Superconducting States, *J. Acous. Soc. Am.*, **28**, pp. 930-944, Sept., 1956.

MAXWELL, J. L.⁵ and FARRAR, H. K.⁵

Automatic Dispatch System for Half-Duplex Teletypewriter Lines, *Commun. and Electronics*, **26**, pp. 441-445, Sept., 1956.

MAYS, J. M.,¹ and BRADY, G. W.¹

Nuclear Magnetic Resonance Absorption by H₂O on TiO₂, *J. Chem. Phys.*, **25**, p. 583, Sept., 1956.

MCLEAN, D. A.¹

Tantalum Capacitors Use Solid Electrolyte, *Electronics*, **29**, pp. 176-177, Oct., 1956.

MESZAR, J.¹

The Full Stature of the Crossbar Tandem Switching System, *Commun. and Electronics*, **26**, pp. 486-496, Sept., 1956.

MOLL, J. L.,¹ TANENBAUM, M.,¹ GOLDEY, J. M.,¹ and HOLONYAK, N.¹

P-N-P-N Transistor Switches, *Proc. I.R.E.*, **44**, pp. 1174-1182, Sept., 1956.

¹ Bell Telephone Laboratories, Inc.

⁵ Pacific Telephone and Telegraph Company, San Francisco, Calif.

MONTGOMERY, H. C.,¹ and BROWN, W. L.¹

Field-Induced Conductivity Changes in Germanium, Phys. Rev., **103**, pp. 865-870, Aug. 15, 1956.

MOORE, E. F.¹

Artificial Living Plants, Sci. Am., **195**, pp. 118-126, Oct., 1956.

MOORE, E. F.,¹ and SHANNON, C. E.¹

Reliable Circuits Using Less Reliable Relays, J. Franklin Institute, **262**, pp. 191-208, Sept., 1956, pp. 281-298, Oct., 1956.

MOSHMAN, J., see Tien, P. K.

NELSON, L. S.¹

Sapphire Lamp for Short Wavelength Photochemistry, J. Am. Optical Soc., **46**, pp. 768-769, Sept., 1956.

NELSON, L. S.,¹ and RAMSAY, D. A.⁸

Absorption Spectra of Free Radicals Produced by Flash Discharges, J. Chem. Phys., Letter to the Editor, **25**, pp. 372-373, Aug., 1956.

NELSON, L. S.,¹ and RAMSAY, D. A.⁸

Flash Photolysis Experiments with a Sapphire Flash Lamp, J. Chem. Phys., Letter to the Editor, **25**, p. 372, Aug., 1956.

NIELSEN, G., JR., see Baldwin, M. W., Jr.

OHM, E. A.¹

A Broad-Band Microwave Circulator, Trans. I.R.E., PGMTT, **MTT-4**, pp. 210-217, Oct., 1956.

OWENS, C. D.

A Survey of the Properties and Applications of Ferrites Below Microwave Frequencies, Proc. I.R.E., **44**, pp. 1234-1248, Oct., 1956.

OWENS, C. D.¹

Modern Magnetic Ferrites and Their Engineering Applications, Trans. I.R.E., PGCP, **CP-3**, pp. 54-62, Sept., 1956.

¹ Bell Telephone Laboratories, Inc.

⁸ National Research Council, Ottawa, Canada.

PHELPS, J. W.¹

Protection Problems in Telephone Distribution Systems, Telephony, **151**, pp. 20-22, 47-48, Sept. 1, 1956.

PRESCOTT, B. E., see Jaycox, E. K.

POWELL, W.⁷

So You've Invested in Modern Plant—Now What?, Telephony, **151**, pp. 27-28, 48-50, Sept., 1956.

RAMSAY, D. A., see Nelson, L. S.

REISS, H.¹

Refined Theory of Ion Pairing. I — Equilibrium Aspects. II — Irreversible Aspects, J. Chem. Phys., **25**, pp. 400-413, Sept., 1956.

REMEIKA, J. P.¹

Growth of Single Crystal Rare Earth Orthoferrites and Related Compounds, Am. Chem. Soc. J., **78**, pp. 4259-4260, Sept. 5, 1956.

RICHARDS, A. P., see Snoke, L. R.

ROBERTSON, S. D.¹

Recent Advances in Finline Circuits, Trans. I.R.E., PGMTT, **MTT-4**, pp. 263-267, Oct., 1956.

SCHMIDT, W. C.¹

Problems in Miniaturization, Proc. Design Engg. Conf., pp. 40-54 May 14-17, 1956.

SEIDEL, H.¹

Anomalous Propagation in Ferrite-Loaded Waveguide, Proc. I.R.E., **44**, pp. 1410-1414, Oct., 1956.

SHANNON, C. E., see Moore, E. F.

SHULMAN, R. G.,¹ and JACCARINO, V.¹

Effects of Superexchange on the Nuclear Magnetic Resonance of MnF_2 , Phys. Rev., Letter to the Editor, **103**, pp. 1126-1127, Aug. 15, 1956.

¹ Bell Telephone Laboratories, Inc.

⁷ New York Telephone Company, New York City.

SHULMAN, R. G.,¹ and WYLUDA, B. J.¹

Nuclear Magnetic Resonance of Si²⁹ in n- and p-Type Silicon, Phys. Rev., Letter to the Editor, **103**, pp. 1127-1129, Aug. 15, 1956.

SLEPIAN, D.¹

A Note on Two Binary Signaling Alphabets, Trans. I.R.E., PGIT, **IT-2**, pp. 84-86, June, 1956.

SNOKE, L. R.,¹ and RICHARDS, A. P.⁴

Marine Borer Attack on Lead Cable Sheath, Sci., Letter to the Editor, **124**, p. 443, Sept. 7, 1956.

SPROUL, P. T.¹

A Video Visual Measuring Set with Sync Pulses, Commun. and Electronics, **26**, pp. 427-432, Sept., 1956.

SUHL, H.¹

The Nonlinear Behavior of Ferrites at High Microwave Signal Levels, Proc. I.R.E., **44**, pp. 1270-1284, Oct., 1956.

TANENBAUM, M., see Moll, J. L.

TIEN, P. K.,¹ and MOSHMAN, J.¹

Monte Carlo Calculation of Noise Near the Potential Minimum of a High-Frequency Diode, J. Appl. Phys., **27**, pp. 1067-1078, Sept., 1956.

THOMAS, D. E.,¹ and DACEY, G. C.¹

Application Aspects of the Germanium Diffused Base Transistor, Trans. I.R.E., PGCT, **CT-3**, pp. 22-25, Mar., 1956.

UHLIR, A., JR.¹

Two-Terminal p-n Junction Devices for Frequency Conversion and Computation, Proc. I.R.E., **44**, pp. 1183-1191, Sept., 1956.

VAN UITERT, L. G.¹

Dielectric Properties of and Conductivity in Ferrites, Proc. I.R.E., **44**, pp. 1294-1303, Oct., 1956.

¹ Bell Telephone Laboratories, Inc.

⁴ Wm. F. Clapp Laboratories, Inc., Duxbury, Mass.

VAN UITERT, L. G.¹

Nickel Copper Ferrites for Microwave Applications, J. Appl. Phys., **27**, pp. 723-727, July, 1956.

WARNER, R. M., JR.,¹ and HITTINGER, W. C.¹

A Developmental Intrinsic-Barrier Transistor, Trans. I.R.E., PGED, **ED-3**, pp. 157-160, July, 1956.

WEINREICH, G.¹

The Transit Time Transistor, J. Appl. Phys., **27**, pp. 1025-1027, Sept., 1956.

WEISS, M. T.¹

Improved Rectangular Waveguide Resonance Isolators, Trans. I.R.E., PGMTT, **Mtt-4**, pp. 240-244, Oct., 1956.

WOLFF, P. A.¹

Theory of Plasma Resonance, Phys. Rev., **103**, pp. 845-850, Aug. 15, 1956.

WYLUDA, B. J., see Shulman, R. G.

YAGER, W. A., see Galt, J. K.

¹ Bell Telephone Laboratories, Inc.

Recent Monographs of Bell System Technical Papers Not Published in This Journal

BALDWIN, M. W., JR., and NIELSEN, G., JR.

Subjective Sharpness of Simulated Color Television Pictures, Monograph 2617.

BOYD, R. C., EBERHART, E. K., HALLENBECK, F. J., PERKINS, E. H., SMITH, D. H., and HOWARD, J. D., JR.

Type-Pl Carrier System — Objectives and Circuit and Equipment Description, Monograph 2644.

CRAWFORD, A. B., and HOGG, D. C.

Measurement of Atmospheric Attenuation at Millimeter Wavelengths, Monograph 2646.

CUTLER, C. C.

The Nature of Power Saturation in Traveling Wave Tubes, Monograph 2647.

EBERHART, E. K., see Boyd, R. C.

FELDER, H. H., PASCARELLA, A. J., SHOFFSTALL, H. F., and LITTLE, E. N.

Intertoll Trunks — Automatic Testing and Maintenance Techniques, Monograph 2652.

FORSTER, J. H., and MILLER, L. E.

Effect of Surface Treatments on Point-Contact Transistor Characteristics, Monograph 2650.

GOHN, G. R.

Fatigue and Its Relation to Mechanical — Metallurgical Properties of Metals, Monograph 2598

* Copies of these monographs may be obtained on request to the Publication Department, Bell Telephone Laboratories, Inc., 463 West Street, New York 14, N. Y. The numbers of the monographs should be given in all requests.

HALLENBECK, F. J., see Boyd, R. C.

HOGG, D. C., see Crawford, A. B.

HOWARD, J. D., JR., see Boyd, R. C.

KELLY, J. L., JR.

A New Interpretation of Information Rate, Monograph 2649.

KING, A. P., and MARCATILI, E. A.

Transmission Loss Due to Resonance of Converted Modes, Monograph 2656.

LINVILL, J. G., and SCHIMPF, L. G.

The Design of Tetrode Transistor Amplifiers, Monograph 2657.

LITTLE, E. N., see Felder, H. H.

MARCATILI, E. A., see King, A. P.

MILLER, L. E., see Forster, J. H.

NIELSEN, G., JR., see Baldwin, M. W., Jr.

PASCARELLA, A. J., see Felder, H. H.

PERKINS, E. H., see Boyd, R. C.

SCHIMPF, L. G., see Linvill, J. G.

SEIDEL, H., see Weisbaum, S.

SHOFFSTALL, H. F., see Felder, H. H.

SMITH, D. H., see Boyd, R. C.

TIEN, P. K.

A Large Signal Theory of Traveling-Wave Amplifiers, Monograph 2610.

TURNER, D. R.

Anode Behavior of Germanium in Aqueous Solutions, Monograph 2570.

VAN HASTE, W.

Statistical Techniques and Electron Tubes for Use in a Transmission System, Monograph 2641.

VAN ROOSBROECK, W.

Theory of the Photomagnetolectric Effect in Semiconductors, Monograph 2607.

WEISBAUM, S., and SEIDEL, H.

The Field Displacement Isolator, Monograph 2661.

Contributors to This Issue

J. F. BAMPTON, B.Sc. in Engineering 1941; British Post Office Engineering Department 1936. Mr. Bampton progressed by competitive examinations to a professional grade in 1942, and in 1944 he was loaned for two years to the government of India to assist with the rapid expansion of carrier telephone and telegraph systems. Since then he has been associated with most of the submerged repeater transmission systems around the British Isles, taking charge of a group in the Transmission and Main Lines Branch of the Engineering Department in 1950. From 1953 to 1956 he was concerned entirely with the transatlantic submarine telephone cable system. Associate Member of The Institution of Electrical Engineers.

M. C. BISKEBORN, B.S. in E.E., S. Dak. School of Mines, 1930; Bell Telephone Laboratories, 1930-42. Western Electric Company, 1942-44; Bell Telephone Laboratories, 1944-. His early work at the Laboratories was concerned with the development of multi-pair carrier and coaxial cables. During World War II, he assisted in the development of one of the first automatic radars, of microwave resonant cavities and of microwave coaxial for the Bureau of Ships. Later, Mr. Biskeborn worked on the development of apparatus for high-frequency electrical measurements on cable. He holds several patents and has written several technical papers including an A.I.E.E. prize paper. At present, he heads a subdepartment on Cable Development. As a part of these responsibilities he was concerned with the design of the transatlantic telephone cable and was responsible for the specifications for it. He is an Associate Member of the A.I.E.E.

F. J. BRAGA, B.E.E., Univ. of Minnesota, 1930; Illinois Bell Telephone Company, 1930-33; Bell Telephone Laboratories, 1934-. Mr. Braga has engaged in development of transmission networks for carrier systems. During World War II he worked on gun computers and networks and circuits for radar applications. He is currently engaged in the development of networks for undersea systems. He is a member of I.R.E.

R. A. BROCKBANK, B.Sc. in Engineering, London University 1922; Ph.D. London University 1934. Dr. Brockbank joined the Research Branch of the British Post Office in 1933 after 10 years in industry, including dielectric research on the original transatlantic cable proposed in 1928. He designed the repeater equipment for the first coaxial cable system in England, 1938. During the war, he was engaged in coaxial developments and on high power negative feedback wideband transmitters. Following the war, he was associated with television transmission over coaxial systems and with submerged repeater development. In 1949 he specialized in this latter work, and since 1953 has been in charge of research and development of submerged repeater systems.

J. W. EMLING, B.S. in E.E., Univ. of Pennsylvania, 1925; Development and Research Department of American Telephone and Telegraph Co., 1925-34; Bell Telephone Laboratories, 1934-. While at A. T. & T. Mr. Emling was particularly concerned with transmission standards and with developing a system of effective transmission rating. He continued this work at Bell Laboratories. In World War II he was concerned with studies in the field of underwater acoustics. Subsequently he has been concerned with systems engineering studies in the fields of engineering economy, voice frequency transmission, rural carrier, radio and television. One of his recent responsibilities covered the early transmission and planning studies of the transatlantic telephone cable system. He is currently Director of Transmission Engineering with responsibility for the systems engineering aspects of exchange and long distance transmission, carrier transmission over wire, telephone stations and some forms of digital transmission. He is a member of the Acoustical Society of America, A.I.E.E., Eta Kappa Nu and Tau Beta Pi.

H. B. FISCHER, B.S. in E.E., Univ. of Wisconsin, 1924; Western Electric Company, 1924-25; Bell Telephone Laboratories, 1925-. Mr. Fischer first worked on broadcast receivers, but after the development of aircraft communication apparatus was started he engaged in the design of aviation receivers. This was followed by work on various types of aviation communication equipment, mobile radio equipment for Bell System use and radio receiving equipment for aircraft instrument landing systems. During the war he worked on various types of electronic equipment for the Armed Forces. Later he worked on overseas radio telephone equipment, video transmission and testing equipment, and submarine communications systems. More recently he has worked on the transatlantic telephone cable project in connection with the manufacture of cable in England.

JOHN M. FRASER, B.E.E., Polytechnic Institute of Brooklyn, 1945; Bell Telephone Laboratories, 1934-. Prior to World War II, Mr. Fraser was concerned with the evaluation of subjective factors affecting the transmission performance of telephone systems. This included the design of equipment for simulating transmission systems in the laboratory. During the war he was chiefly concerned with the design and evaluation of communication systems for the military. Later he was engaged in transmission work on long distance carrier systems. On the transatlantic telephone cable he was mainly concerned with the System Engineering aspects of the cable system. He is a member of Sigma Xi, Tau Beta Pi and Eta Kappa Nu.

T. F. GLEICHMANN, B.E., Johns Hopkins Univ., 1929; Bell Telephone Laboratories, 1929-. Until 1942 Mr. Gleichmann was engaged in the design and development of open wire and cable carrier systems. During the period of World War II he worked on the design and development of radio telephone pulse communication systems for military and Bell System applications. He then engaged in development work in connection with coaxial cable carrier systems. He was in charge of the group responsible for the circuit and mechanical design of the repeater unit for the transatlantic submarine telephone cable. He is a member of Tau Beta Pi.

R. G. GRIFFITH, graduate I.E.E., London 1924. Studied general engineering in Royal Naval Air Service and Communication Engineering in London. Mr. Griffith left England in 1924 to join All-American Cables Inc. (now American Radio and Cable Corporation) becoming Project Engineer in 1925, supervising ac telegraph transmission superimposed tests (then termed "wired wireless") on dc duplex telegraph cable between Balboa Canal zone and Fishermans Point, Cuba. He developed and supervised the introduction of the synchronous fork cable signal regenerator, which established the through cable circuits between New York and Buenos Aires via the west coast cables of South America. In 1929 Mr. Griffith was appointed Assistant Chief Engineer of Creed and Company, and in 1932 was placed in charge of development. In 1935 he joined Cable and Wireless Limited. From 1943 to 1946 he was loaned to the foreign office communication center in charge of special machine cipher development. He became Chief Engineer of Cable and Wireless London Communications center in 1946, and in May 1954 joined the Canadian Overseas Telecommunication Corporation as Chief Engineer. Mr. Griffith holds some 60 patents relating to telecommunications.

R. J. HALSEY, B.Sc. in Engineering, London University, City and Guilds College, 1926; Diploma of the Imperial College, 1926. Mr. Halsey entered the Engineering Research Branch of the British Post Office in 1927 where he was engaged on line transmission problems including, from 1938, the design of submerged repeaters and systems. In 1947 he became Head of the Line Transmission Division, and in 1952, Assistant Engineer-in-Chief concerned with all submarine cable matters; in this capacity, his primary concern has been the transatlantic submarine telephone cable. Associate of the City and Guilds of London Institute and Member of the Institution of Electrical Engineers.

WILLIAM W. HEFFNER, B.S. in Industrial Engineering, Pennsylvania State University 1929; Western Electric Company, Kearny, New Jersey, 1929-1932; Consulting work on industrial engineering in the Management Field 1932-1936; Western Electric Company 1936-. Mr. Heffner's initial work at Western was concerned with jacks, keys, and mica capacitors. In 1942, he became a Department Chief in charge of manual telephone apparatus. In 1947 he was made an Assistant Superintendent in engineering for manual apparatus. His assignments continued through 1952 in engineering for several manufacturing engineering functions, including factory engineering, manufacture of manual apparatus and equipment, metal finishing, material handling, and packing. Between 1952 and 1954 he was Assistant Superintendent in charge of the Relay Assembly Shops at Kearny, New Jersey. In 1954 he was placed in charge of operating, production control, plant operations, and maintenance at Hillside, New Jersey, where the flexible repeaters for the transatlantic submarine telephone cable were manufactured. More recently, he was placed in charge of the Fairlawn, New Jersey, shop of Western Electric, where telephone apparatus and switching equipment are being built. Mr. Heffner is a member of Sigma Tau.

M. F. HOLMES, B.Sc. in Physics 1937; British Post Office 1938-. Mr. Holmes transferred to the Engineering Department in 1942 and since 1944 has been concerned primarily with thermionics. He is now engaged in the study of factors leading to changes of tube characteristics.

JOHN S. JACK, Mountain States Telephone and Telegraph Company 1919-1930; American Telephone and Telegraph Company, Long Lines Department, 1930-. Mr. Jack was engaged in various Plant assignments in Colorado and Wyoming between 1919 and 1930. In 1930, he became Division Outside Plant Engineer for Long Lines in Denver; in 1938 he

transferred to Chicago as Division Plant Engineer, and three years later, moved to Omaha, Nebraska as District Plant Superintendent. He was transferred to the Personnel Department in New York in 1945 as General Supervisor of Wages and Working Practices. In 1949 he returned to the Plant Department as General Construction Supervisor, and in 1951 became Engineer of Outside Plant. In 1953 he was appointed Assistant General Manager — Special Projects; in this capacity he helped direct construction of the transatlantic submarine telephone cable. Mr. Jack is a licensed professional engineer in Nebraska.

M. J. KELLY, B.S., Missouri School of Mines and Metallurgy, 1914; M.S., Univ. of Kentucky, 1915; Ph.D., Univ. of Chicago, 1918; honorary degrees — D.Eng., Univ. of Missouri, 1936; D.Sc. Univ. of Kentucky, 1946; LL.D., Univ. of Pennsylvania, 1954; D.Eng., New York Univ., 1955; D.Eng., Polytechnic Institute of Brooklyn, 1955. Western Electric Company, 1918–25. Bell Telephone Laboratories, 1925–. Dr. Kelly became Director of Vacuum Tube Development in 1928; Development Director of Transmission Instruments and Electronics, 1934; Director of Research, 1936; Executive Vice President, 1944; President, 1951. He was awarded the Presidential Certificate of Merit in recognition of his contributions in World War II and now serves on several advisory boards in the Department of Defense and the Department of Commerce. Dr. Kelly is a Fellow of the American Physical Society, the Acoustical Society of America, I.R.E., and A.I.E.E. He is a Foreign Member of the Swedish Royal Academy of Sciences and a member of the National Academy of Sciences, the American Philosophical Society, Sigma Xi, Tau Beta Pi and Eta Kappa Nu. He is a Life Member of the M.I.T. Corporation and a Trustee of Stevens Institute of Technology. His honors include the Air Force Association Trophy in 1953; the Industrial Research Institute Medal in 1954; and the Christopher Columbus International Communication Prize in 1955.

R. KELLY, Associate of Royal College of Science, Ireland 1925; B.Sc. University College, Dublin 1937. After four years experience on power work for the Dublin United Tramways, he joined the power section of Standard Telephones and Cables in 1925. Following ten years of laboratory and field experience on carrier telephone and VF telegraph equipment, he took charge in 1936 of power development for transmission equipments.

HAROLD A. LAMB joined the Western Electric Company Installation Department in 1920, where he became engaged in installation and

installation engineering of telephone equipment. In 1923, he entered the Engineer of Manufacture Organization at Hawthorne, where he became concerned with relays, panel, step-by-step, and crossbar machine switching apparatus. During this time, he attended the Lewis Institute of Technology. In 1936 he was appointed a Department Chief on step-by-step apparatus. During World War II, Mr. Lamb transferred to the Passaic, New Jersey, Shops as Assistant Superintendent in the Western Electric Radio Division. Here he was engaged in engineering the manufacture of submarine radar and radar bomb sights. Returning to the Western Electric, Kearny, New Jersey, Works in 1947, he was concerned principally with central office apparatus, including the card translator, and in 1953 was placed in charge of the Hillside, New Jersey, Engineering and Inspection Organizations for building the flexible repeaters for the transatlantic submarine telephone cable. He is at present Resident Head of the Hillside Shops on flexible repeater manufacture.

ANDREW W. LEBERT, B.S. in E.E., New York Univ., 1932; Cornell-Dubilier Corporation, 1932-1936; Bell Telephone Laboratories, 1936-. For the first five years at the Laboratories, Mr. Lebert worked on transmission engineering on open wire and cable carrier systems. He then was concerned with fault location problems. During World War II, he turned to military communications on cable and open wire, and, following this period, he spent eight years on coaxial cable systems development. Since 1952 he has been connected with transatlantic telephone cable development. He is a member of I.R.E., Tau Beta Pi and Psi Upsilon.

CAPT. W. H. LEECH entered the British Post Office in 1920 as Third Officer of H.M.T.S. *Alert* and was later promoted to the old H.M.T.S. *Monarch* of which ship he became Chief Officer. Both of these ships were subsequently lost by enemy action in World War II. After a year ashore as Assistant Submarine Superintendent in 1938-39 he took command of H.M.T.S. *Aeriel* and, in 1940, of H.M.T.S. *Iris*. In 1944 his ship was engaged in laying cables to the Normandy Beach head, an operation for which he was awarded the Distinguished Service Cross. In 1946 he became Submarine Superintendent, in immediate charge of the Post Office cable fleet and as such, directed the operations of H.M.T.S. *Monarch* during the laying of the transatlantic cables. He is an Officer of the Order of the British Empire (O.B.E.).

HERBERT A. LEWIS, E.E., Cornell Univ., 1926; Bell Telephone Laboratories, 1926-. Before World War II Mr. Lewis worked on the design

of equipment for manual and dial central offices, PBX's and broad-band carrier installations. During the war, he was concerned with the mechanical design of radar systems for the military. He was later responsible for transmission and equipment development for various carrier telephone systems. As project engineer for the Laboratories phases of the transatlantic telephone cable he was responsible for its transmission and equipment development. He is now Director of Outside Plant Development and is responsible for the devising and developing of new and improved methods, materials and equipment for that part of the telephone network which connects one central office with another and which ties the telephone customer's equipment into the central office. He is a senior member of I.R.E.

ARTHUR H. LINCE, B.S. in E.E., Univ. of Michigan, 1925; Bell Telephone Laboratories, 1925-. Until 1941 Mr. Lince worked on the engineering and design of dial central office equipment. During World War II he was concerned with engineering and design of radar for the armed forces. He then became involved with the development of microwave antennas, towers, waveguides and related items for microwave radio relay systems and testing equipment. He was engaged in the building of repeaters for the Havana-Key West submarine cable. Beginning in 1953, he has been in charge of the group responsible for the design of water-tight enclosures for the repeaters used on the transatlantic telephone cable.

G. H. LOVELL, B.S. in E.E., Texas A & M College, 1927; M.S. in E.E., Polytechnic Institute of Brooklyn, 1943; N.Y. Edison Co., 1927-28; Bell Telephone Laboratories, 1929-. From 1929 until 1948 Mr. Lovell was concerned with the development of crystal filters for carrier systems. He then worked on the development of networks for use in broad-band amplifiers. For the transatlantic telephone cable project he worked on the amplifier networks and the equalization of the undersea system.

J. O. McNALLY, B.S. in E.E., Univ. of New Brunswick, Canada, 1924; Western Electric Company, 1924-25; Bell Telephone Laboratories, 1925-. Mr. McNally has specialized in research and development on electron tubes for Bell System communication and military uses. This included work on voice and carrier repeater tubes, electron tubes for the first commercial transatlantic radio system, and for the talking movie industry. During World War II, he had development responsibility for

many of the klystrons used in radar equipment. Later he again became concerned with the development of long-life tubes for submarine telephone cables. He has been awarded several patents on electron tube construction and operation. He is a Fellow of the I.R.E. and a member of the American Physical Society.

GEORGE W. MESZAROS, B.E.E. 1939, College of the City of New York; Bell Telephone Laboratories, 1926-. Mr. Meszaros started his Bell System career in the Systems Drafting Department. After spending a short time in several engineering groups of the System Department, he transferred to the Power Development Department in 1941. Here he has specialized in electronically controlled power equipment. Currently he is in charge of a group designing transistorized power supplies for the electronic switching system and for several military projects.

G. H. METSON, B.Sc. in Engineering, University of London 1931; M.Sc. in 1938 and Ph.D. in Applied Science and Technology, Queens University, Belfast 1941. Dr. Metson is in charge of the Thermionics Group at the Post Office Research Station and is particularly concerned with oxide coated cathodes and problems of tube life. He was responsible for the tubes used in the British submerged repeaters. Member of the Royal Institution and an Associate Member of The Institution of Electrical Engineers.

ELLIOTT T. MOTTRAM, B.S. Columbia University 1927, M.E. 1928; Western Electric Company 1922-25; Bell Telephone Laboratories, 1928-. Mr. Mottram's first assignments were in the development of disc recording and reproducing machines and equipment. Later he was concerned with sound on film recording and reproducing equipment, and with tape recording. From 1939 to 1950, he was engaged in development of airborne radio and radar equipment, electronic computer and bomb sights, and airborne homing missiles. As Director of Transmission Systems Development since 1950, he has been concerned with the development of transmission systems and equipment for military purposes, transmission test equipment, and television and wire transmission systems. In this capacity, he was responsible for technical liaison with the British Post Office on submarine cable matters and was in charge of Laboratories' activities in this field. He is a member of the A.S.M.E. and I.R.E.

SIR GORDON RADLEY, B.Sc. in Engineering, University of London 1919; Ph.D. University of London 1934. Sir Gordon's under-

graduate studies were interrupted by military service in World War I. Engineering Research Branch of the British Post Office 1920, where he was engaged initially on materials problems and later on interference, corrosion, and long distance signaling. He became Head of Research in 1939, and in 1949 was made Deputy Engineer-in-Chief. In 1951 he became Engineer-in-Chief, and in this position, was one of the principal architects of the transatlantic submarine telephone cable system. In 1954 he was made Deputy Director General, and in 1955 Director General — the permanent Head of the British Post Office. He became a Commander of the Order of the British Empire (CBE) in 1946 and was honored with a knighthood in 1954. In 1956, he became a Knight Commander of the Bath (KCB) and is President of The Institution of Electrical Engineers for the year 1956-57.

H. H. SPENCER, B.S. in M.E., Univ. of New Hampshire, 1923; Bell Telephone Laboratories, 1923-. He has been engaged primarily in the development of power supplies for broadband carrier, long distance and repeater equipment, including automatic plants for unattended operation on J, K, and L carrier systems and TD-2 microwave radio relay systems. Mr. Spencer is an Associate Member of the American Institute of Electrical Engineers.

J. F. P. THOMAS, B.Sc. London University 1942; British Post Office Research Branch, 1937-. In his early years, Mr. Thomas was engaged on investigations into contact phenomena and dust core magnetic materials. In 1948, he was transferred to the Submerged Repeater Group, where his main work has been the design and construction of power feeding equipment and pulse monitoring equipment used for fault location in submerged repeater systems. Associate Member of The Institution of Electrical Engineers.

REXFORD S. TUCKER, A.B., Harvard College, 1918; S.B., Harvard Engineering School, 1922; American Telephone and Telegraph Company, 1923-34; Bell Telephone Laboratories, 1934-. Mr. Tucker's early work was on noise and crosstalk prevention. During World War II he was engaged in classified military projects and served as co-editor of a War Department technical manual, *Electrical Communications Systems Engineering*. After the war he worked on mobile radio systems engineering and then the transatlantic telephone cable. He is an Associate Member of A.I.E.E., Senior Member of I.R.E., Charter Member of Acoustical Society of America, member of Sub-Committee No. 1 of

American Standards Association Sectional Committee C63, Harvard Engineering Society, and Phi Beta Kappa.

EDMUND A. VEAZIE, B.A. in Physics, Univ. of Oregon, 1927; Bell Telephone Laboratories, 1927-. His early assignments included the design of multi-grid tubes for use in aircraft radio receivers, police transmitters, and carrier telephone systems. During World War II he concentrated on tubes for radar and other military applications, including proximity fuses and gun directors. Since then he has been engaged principally in the design, fabrication control, testing, and selection of tubes for use in submarine telephone cable systems. He holds several patents on electron tubes and associated circuits. He is a Senior Member of I.R.E. and a member of Phi Beta Kappa.

D. C. WALKER, B.Sc. in Engineering and Diploma of the Imperial College from the City and Guilds College, University of London, 1937; British Post Office Research Branch 1938. Mr. Walker's early work was on interference and protection problems and during the war on special investigations for the services. Later engaged on development and equipment for carrier telephone systems, and since 1946 has specialized on submerged repeater systems. He is in charge of the group concerned with the design of the internal electrical unit of the rigid transatlantic telephone cable repeaters and the special terminal equipment. Associate Member of The Institution of Electrical Engineers.

V. G. WELSBY, B.Sc. London University 1934; Ph.D. London University 1946; Research Branch of the British Post Office, 1936. Dr. Welsby was at first a member of a group dealing with the design of multichannel carrier apparatus. Since 1947 he has been engaged in submerged repeater development, and during the last few years, has been in charge of the group concerned with the mechanical design of repeater housings and glands, and with the laying of rigid repeater systems. His Ph.D. degree was awarded for his work on dust-cored inductors. He is the author of a text book on inductor theory and design. Associate Member of The Institution of Electrical Engineers.

M. C. WOOLEY, B.S. in E.E., Ohio Northern Univ., 1929; Bell Telephone Laboratories, 1929-. Mr. Wooley was engaged in the development and design of inductors until 1935. Capacitor development then occupied his attention until 1949, concluding with the development and produc-

tion of capacitors for the Key West-Havana submarine cable repeaters. He then became concerned with fundamental development, primarily on materials and processes used in capacitors, including those for the transatlantic submarine telephone cable. He is currently supervising a group engaged in development and design of capacitors and resistors for submarine cable repeater applications for other systems. He is a member of Nu Theta Kappa.

2347
H E B E L L S Y S T E M

Technical Journal

VOTED TO THE SCIENTIFIC AND ENGINEERING
PECTS OF ELECTRICAL COMMUNICATION

LUME XXXVI

MARCH 1957

KANSAS CITY NUMBER 2

A New Carrier System for Rural Service

APR 3 1957

R. C. BOYD, J. D. HOWARD, AND L. PEDERSEN 349

An Experimental Dual Polarization Antenna Feed for Three Radio
Relay Bands

R. W. DAWSON 391

The Character of Waveguide Modes in Gyromagnetic Media

H. SEIDEL 409

Measurement of Dielectric and Magnetic Properties of Ferro-
magnetic Materials at Microwave Frequencies

WILHELM VON AULOCK AND JOHN H. ROWEN 427

Sensitivity Considerations in Microwave Paramagnetic Resonance
Absorption Techniques

G. FEHER 449

The Determination of Pressure Coefficients of Capacitance for
Certain Geometries

D. W. McCALL 485

Reading Rates and the Information Rate of a Human Channel

J. R. PIERCE AND J. E. KARLIN 497

Binary Block Coding

S. P. LLOYD 517

Selecting the Best One of Several Binomial Populations

MILTON SOBEL AND MARILYN J. HUYETT 537

Bell System Technical Papers Not Published in This Journal 577

Recent Bell System Monographs 583

Contributors to This Issue 588

THE BELL SYSTEM TECHNICAL JOURNAL

ADVISORY BOARD

A. B. GOETZE, *President, Western Electric Company*

M. J. KELLY, *President, Bell Telephone Laboratories*

E. J. McNEELY, *Executive Vice President, American Telephone and Telegraph Company*

EDITORIAL COMMITTEE

B. McMILLAN, *Chairman*

S. E. BRILLHART

A. J. BUSCH

L. R. COOK

A. C. DICKIESON

R. L. DIETZOLD

K. E. GOULD

E. I. GREEN

R. K. HONAMAN

H. R. HUNTLEY

F. R. LACK

J. R. PIERCE

G. N. THAYER

EDITORIAL STAFF

J. D. TEBO, *Editor*

R. L. SHEPHERD, *Production Editor*

T. N. POPE, *Circulation Manager*

THE BELL SYSTEM TECHNICAL JOURNAL is published six times a year by the American Telephone and Telegraph Company, 195 Broadway, New York 7, N. Y. F. R. Kappel, President; S. Whitney Landon, Secretary; John J. Scanlon, Treasurer. Subscriptions are accepted at \$5.00 per year. Single copies \$1.25 each. Foreign postage is 65 cents per year or 11 cents per copy. Printed in U. S. A.

THE BELL SYSTEM TECHNICAL JOURNAL

VOLUME XXXVI

MARCH 1957

NUMBER 2

Copyright 1957, American Telephone and Telegraph Company

A New Carrier System for Rural Service

By R. C. BOYD, J. D. HOWARD, JR, and L. PEDERSEN

(Manuscript received July 19, 1956)

A study of the problem of providing telephone service to rural customers indicated the need for a flexible carrier system that could be used economically on new and existing rural cable and open wire lines. The desire for low cost required new approaches to almost every phase of the carrier system design for rural service, which has been designated Type P1.

Use of transistors led to sweeping changes in the detailed circuitry and also created demand for other new components. Mounting and interconnecting the circuit components by means of printed wiring boards emphasized the necessity for close coordination between design and manufacturing objectives. The low power-drain requirements of transistor circuitry were supplied economically by the use of similar solid state devices, a new storage battery, and efficient packaging.

A fast, accurate and simple method has been evolved for applying the P1 carrier system to rural lines with a minimum of line treatment or rearrangement. Plug-in equipment, readily accessible test points, and a carrier test set provide the ease of maintenance needed in the use of telephone equipment at remote locations. Use of the P1 carrier system will extend the application of electronic equipment outside of the telephone central office and provide a carrier system whose performance will be consistent with requirements for high quality communication service at low cost.

1. INTRODUCTION

Although carrier has been used successfully to provide trunks in the Bell System for more than 35 years, it has not been economically feasible, up to the present time, to apply carrier telephone techniques extensively to the rural telephone plant. The technical and economic problems associated with providing telephone service to customers in rural areas has long been one of the most difficult problems facing the telephone industry. The widely scattered locations of customers in rural areas have led to a large number of rural telephone routes with only a few customer lines per route. This has precluded the use of large cables on any one route, which would be economically attractive in urban areas. The extensive use of carrier has not been feasible because the distances from the rural customers to the Central Office are in the 5- to 20-mile range in which carrier has not been generally economical in the past.

The two lines of attack which were taken on this problem were to reduce the cost of telephone plant through less expensive small cables and open-wire plant,¹ and to provide an economically attractive carrier system designed to meet the particular needs of rural telephone service.

This paper discusses the broad objectives for a rural customer carrier system, the major parameters of the P1 system which was developed to meet those objectives, and its circuit, equipment, and power arrangements. It also covers the engineering and maintenance methods to be used by the Bell System Operating Companies to install and operate the system.²

2. BROAD OBJECTIVES FOR P1 CARRIER SYSTEM

The broad objectives for the Type P1 carrier system resulted from the stringent economic limits imposed on the system to enable it to prove in over conventional rural plant of the latest and most economical design, from the requirements of rural telephone transmission and signaling, and from Bell System experience gained with earlier carrier systems for customer and trunk use. The low cost objective for this system also implied the need to achieve an appropriate balance among an economic first cost of equipment, low in-place cost due to simplified engineering and installation practices, and accompanying low annual costs due in part to simplified system maintenance.

To achieve an economic carrier system for rural telephone use, the dc power requirements of the terminals and repeaters had to be kept low.

¹ Lester Hochgraf and R. G. Watling, Telephone Lines for Rural Subscriber Service, A.I.E.E. Communication and Electronics, No. 18, p. 171, May, 1955.

² These aspects of the P1 system are covered in more detail in four papers on "The P1 Carrier System." A.I.E.E. Communication and Electronics, No. 24, pp. 188, 191, 195, 205, May, 1956.

This was especially important since previous Bell System experience indicated that where commercial power is used to supply the system, some form of reserve must be provided, and where commercial power is not available, the use of primary batteries places a premium on minimizing the power required.

From these considerations two additional major objectives were derived: low manufacturing costs for the components and assembled equipment, and the use of transistors to minimize power supply drains. In addition, flexibility was needed in the proposed carrier system because of the difficulty of accurately forecasting the demand for rural service.

These objectives have been met in the design of the Type P1 rural customer telephone system. It is a fully-transistorized system consisting of independent two-way carrier channels applicable in increments of one to four at a time in the frequency band above the regular voice frequency circuit. Each channel uses a terminal at the central office and at a remote point with intermediate repeaters as necessary. Between terminals, the system is equivalent to a rural voice frequency line with no changes required in the central office or rural customer equipment. Beyond the outlying terminal, distribution is by voice frequency wire on a single or multiparty basis. The system can be applied to existing and new lines utilizing combinations of fine gauge exchange cable and copper or steel open-wire. Systems can be used on each of several pairs on a given pole line, the number depending on the line characteristics.

3. MAJOR PARAMETERS OF P1 CARRIER SYSTEM

This section summarizes the important features incorporated in the P1 carrier system and the reasons governing their choice. The system has a number of features in common with Bell System toll carrier systems, but it also differs in several important aspects because of specific rural requirements. One aspect is the signaling, which requires different arrangements at the two ends of the circuit because of the widely different signals carried in the two directions. Another is that the remote terminals of the individual channels are usually distributed along the line rather than grouped at a common location.

3.1 *Transmission Plan*

It is difficult to divorce the considerations leading to the choice of carrier frequency range from those affecting the choice of modulation in the carrier system. Studies of growth on rural lines indicated that a system giving three or four channels (customer circuits) on one pair of wires, in addition to the physical circuit, should be sufficient if systems could be applied to each of several pairs on a given open-wire line.

The blocking out of the frequency range was controlled by a number of factors. Cost considerations required that the carrier frequencies be kept above the voice frequency range. If carrier extended into the voice range, the voice frequency circuit would be lost on a carrier pair. One of the carrier channels applied to that pair would have to be used to replace it. Thus, the addition of four carrier channels to a pair would yield a net gain of only three channels. This in turn would increase the net cost per gained channel. Filter costs determined how close to the voice frequency band the carrier frequency range could be placed and in conjunction with the number of channels required how closely the channels could be placed to each other.

Crosstalk considerations restricted the carrier frequency range to below about 100 kc in order to reduce the cost of line treatment and rearrangement of pairs on existing rural lines. By using this frequency range it appeared possible to apply more than one carrier system to crossarms on an open-wire route. The rapid rise in attenuation with frequency of steel wire used on rural lines dictated that the range of frequencies be kept low. As a result of these two sets of considerations, development work on the P1 carrier system was concentrated in the 8- to 100-kc range.

Amplitude modulation of the carrier frequencies was chosen over other forms of modulation because of the simpler terminal circuitry and equipment and because of the saving in bandwidth. Use of amplitude modulation and the use of compandors, discussed in a later section, were felt to compensate for possible transmission advantages that could

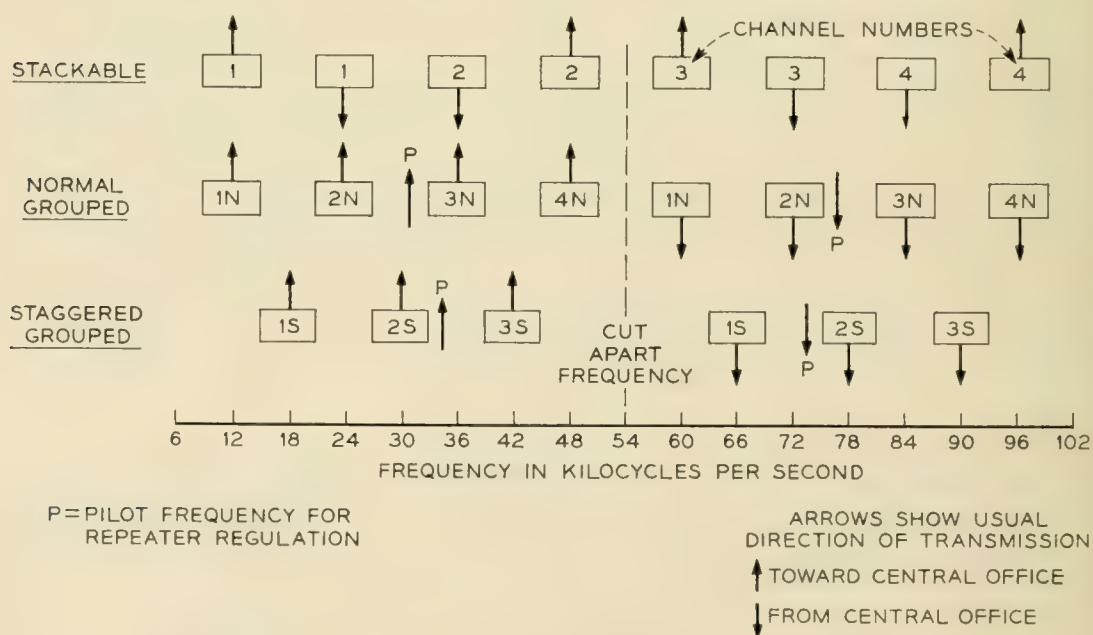


Fig. 1 — Type P1 Carrier Frequency Plan.

be obtained by using angular modulation (frequency or phase) with a large modulation coefficient.

Cost was again a major factor in the choice between double sideband and single sideband amplitude modulation. Past experience with other carrier systems has indicated that filters are a major part of the cost of a system and, when frequency space is available, double sideband filters are, in general, less expensive than those for single sideband. In addition, the cost of a single sideband system would be increased because of the problem of obtaining the necessary carrier supply at the terminal.

The frequency plan developed for the P1 carrier system is shown in Fig. 1. The unusually wide carrier spacing of 12 kc was adopted in order to minimize filter costs. Since the remote terminals are generally distributed along the line, it was not practical to use double modulation to accomplish filtering in the most efficient frequency range. Instead, filtering was done at line frequencies. Every effort was made to achieve channel filter designs with maximum efficiency of element utilization. Advantage was taken of the more leisurely rising characteristics of the double sideband filters permitted by the wide frequency spacing.

The stackable frequency arrangement was provided for non-repeated operation, because when the lowest two carrier frequencies are used to provide a channel, it can be used over substantially longer distances than channels using higher frequencies. The grouped arrangements were provided for repeated systems to reduce the cost and number of the repeater filters and amplifiers needed to separate the two directions of transmission. The staggered grouped arrangement can be used with the normal grouped arrangement on a pole line having poor crosstalk coupling in order to increase the effective coupling loss between carrier channels on different pairs. The grouped and stackable arrangements cannot be used on the same pole line, because certain frequencies would be used for both directions of transmission. This would produce large differences between transmitted and received carrier power at terminals and repeaters which would lead to intolerable crosstalk.

A number of terminal arrangements were studied in order to implement the above frequency plan. The arrangement for a remote terminal shown in Fig. 2 was chosen as the simplest terminal meeting all of the system requirements. It is very similar to the channel terminal arrangement used in the Type N1 carrier system, another double sideband amplitude modulation system used for long distance trunks of the Bell System. The several shaded portions in the figure show the breakdown of the terminal functions into individual sub-units, which are the basis for the equipment arrangements discussed in Section 5 of this paper. A number of the other important features that make up the terminal arrangement are discussed in the following sections.

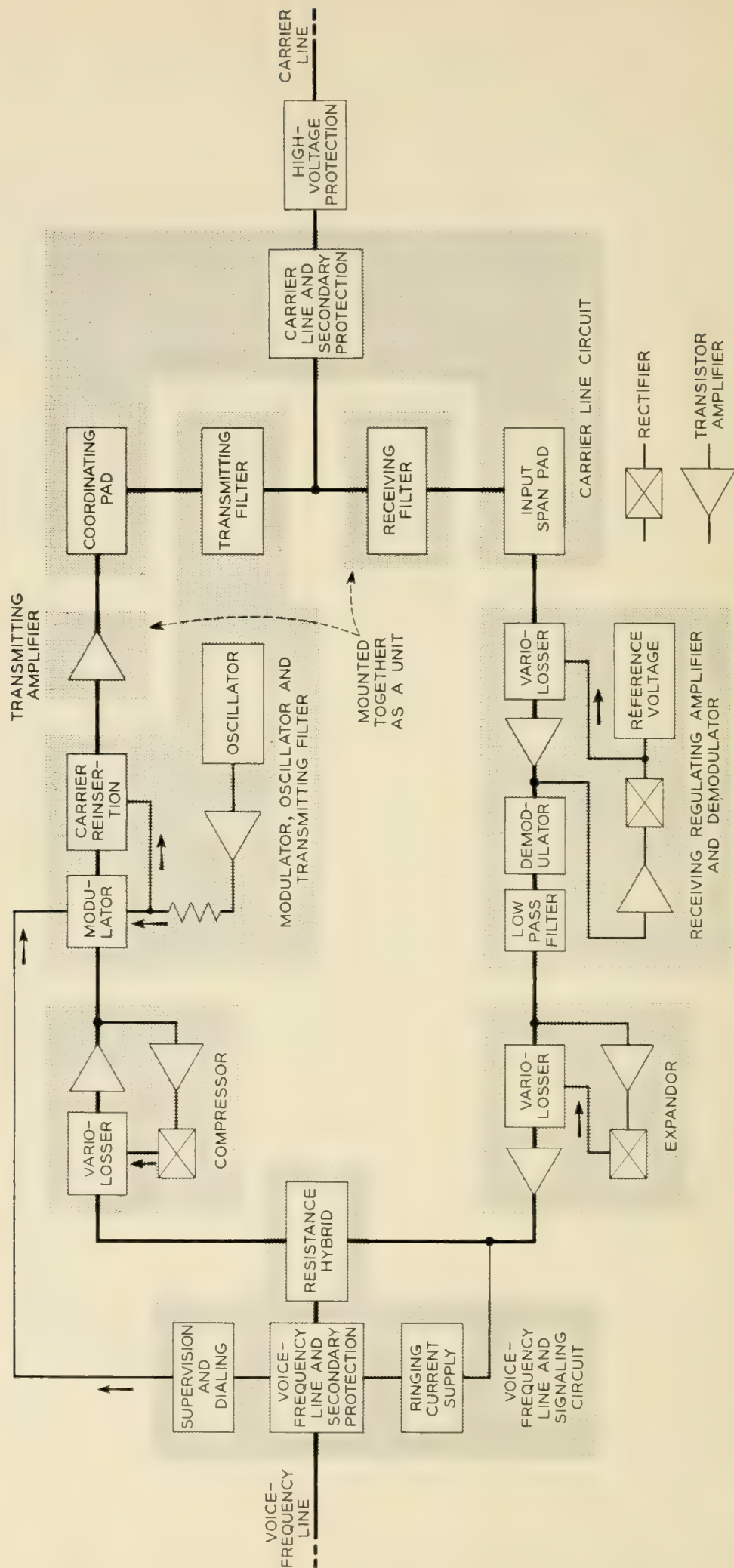


Fig. 2 -- P1 Carrier remote terminal block diagram.

3.2 Use of Transistors

Transistors were chosen for use in the P1 system because they are low voltage, low power devices as compared to electron tubes suitable for transmission circuitry. Also, transistors are expected to be lower in cost and inherently longer life devices than electron tubes, thus contributing to reduced initial and operating costs.

The dc power requirements for the P1 system, using transistors, may be compared to those for a channel terminal in the Type N1 system as an indication of the dc power saving that has been achieved with the P1 system. A transistorized P1 terminal requires about 1.2 watts while it is in operation compared to 40 watts required for an N1 terminal,

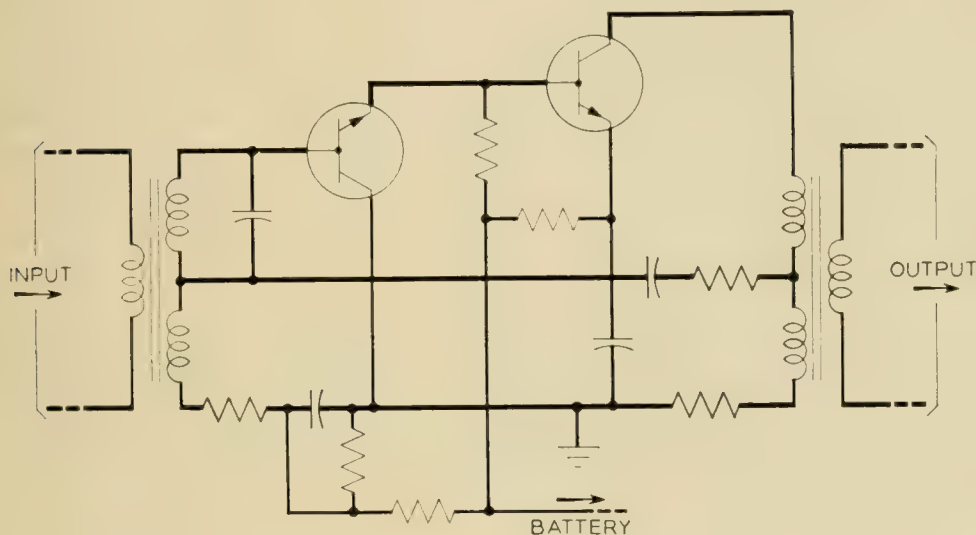


Fig. 3. — P1 transistor transmitting amplifier circuit.

which represents a substantial power reduction achieved by the use of transistors. Because part of the P1 terminal is turned off during idle periods, the average power required over a day is about 0.9 watt.

During the development of the P1 terminals it was found that a single design of a transistor amplifier could be used in several different places. These included the compressor and expander amplifiers, the transmitting amplifier and the input portion of the receiving amplifier. The circuit for that amplifier is shown in Fig. 3.

The amplifier uses Western Electric NPN grown junction type transistors coded 4B for the voice frequency amplifiers and 4C transistors for the carrier amplifiers. The first transistor is connected as a common collector and the second as a common emitter. By using them in this manner it is possible to employ the same type of transistor in both stages. Feedback is obtained by using hybrid coils at both the input and output

of the circuit in much the same manner as for electron tube circuits. One significant difference is that in this transistor circuit only the second transistor introduces a 180-degree phase shift. This permits both input and output coils to return to a common ground as in a three-tube electron-tube circuit and thereby avoids the circuit complications of a two-tube circuit where one of the coils must float off ground. A simple resistance interstage is used and battery filtering completes the circuit.

3.3 Low Voltage Protection

Use of transistors gave rise to the need for supplementary protection from voltage surges on the line below those for which conventional carbon blocks afford protection. This additional protection was obtained by using the reverse voltage breakdown characteristics of newly developed silicon-aluminum junction diodes as shown in Fig. 4. Protection is provided in the 50- to 1,000-volt range by the diodes and above a nominal value of 750 volts by the carbon blocks. During the normally short period of operation the small diodes carry a current of up to 10 amperes.

3.4 System Levels and Carrier Line Loss

The carrier frequency output power of the transistorized transmitting amplifier in the terminals was set at +6 dbm. This level was limited primarily by the power handling capabilities of the transistors used. Because of the loss of the secondary protection circuitry, band filters, and the line transformer, this became +4 dbm at the carrier line terminals. This is equal to the highest carrier power transmitted by the Type N1 carrier system. With 50 per cent modulation of the carrier, the effective sideband level at the transmitting line terminals is only 2 db below that transmitted by the Type O carrier system, the most recent carrier system used for open-wire long distance trunks of the Bell System.

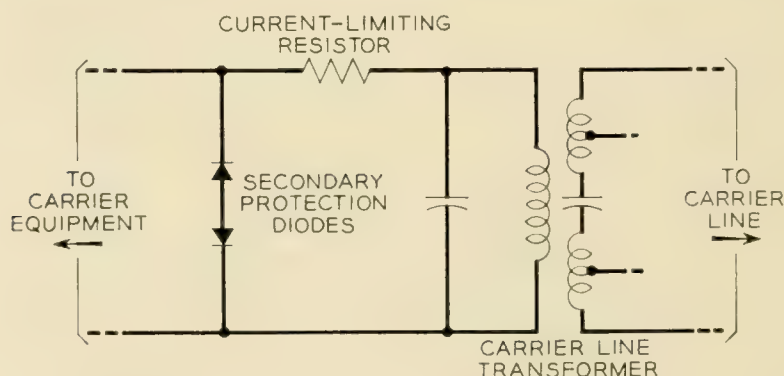


Fig. 4 — Secondary protection circuit in carrier portion of P1 carrier terminal.

The +4 dbm output level coupled with noise and crosstalk considerations indicated that 30-db bare carrier line loss between terminals would be possible. A survey of existing and planned Bell System rural telephone lines indicated that substantial amounts of entrance cable and open wire would be encountered in potential carrier layouts. Calculations of carrier frequency loss of those facilities showed that 30-db loss would not be sufficient to care for all of the necessary rural applications, which confirmed the need for carrier repeaters.

3.5 *Compandors*

Compandors were incorporated in the P1 system because their several advantages more than offset their added cost. The crosstalk and noise advantage provided by their use reduced the need for expensive line treatment to reduce crosstalk. In addition, the compandor noise advantage permitted lower received carrier levels to be used, thus increasing the permissible carrier line loss. Compandors also eased the requirements on terminal and repeater filters, thus reducing filter cost.

The compandor in the P1 system is a simplified version of the syllabic compandor used in Type N1 and O carrier systems, but its performance is comparable to those units. The new problem of matching the compressor and expander characteristics in P1 terminals operating in different ambient temperatures has been simplified by the use of silicon-aluminum junction diodes in the compandor variolossers and control circuits.

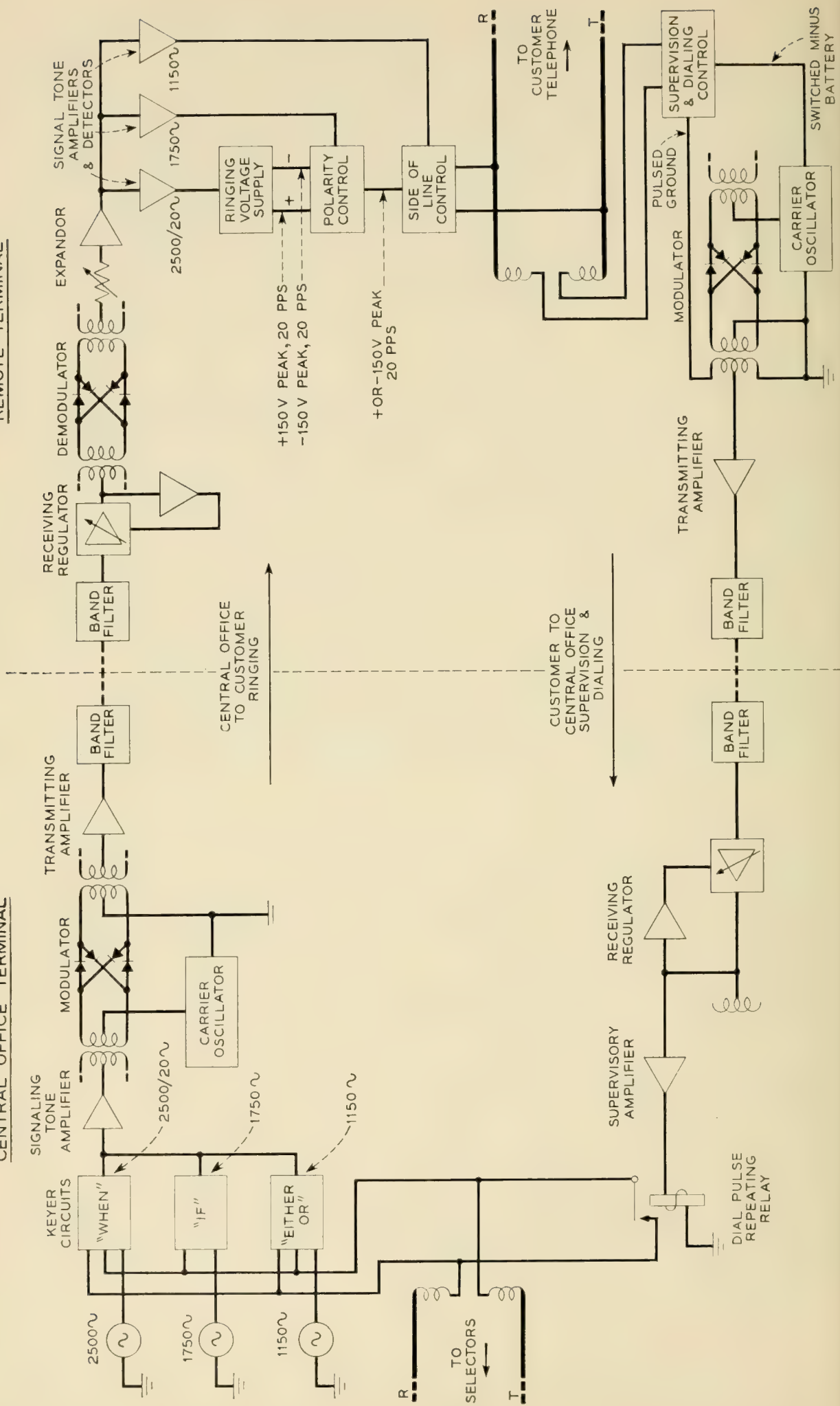
3.6 *Channel Regulation*

Channel regulation was necessary to provide satisfactory transmission performance and keep maintenance adjustments to a minimum. The regulation was designed to compensate for daily and seasonal carrier circuit net loss variations caused by changes in line attenuation with temperature. It would be desirable to have the terminal regulation range equal to 30 db, the maximum line loss that can be spanned between the terminals, to ease engineering layout considerations. However, cost considerations led to a 15-db range, with span pads used where required by system layout to adjust the received carrier power to the center of the range of the regulator.

The regulation in the receiving amplifier is of the backward-acting type. A reference control signal, derived from the receiving amplifier output, is used to vary the loss of the balanced diode variolossers at the amplifier input. The regulator "stiffness" of 1.6-db change in channel voice frequency output for 15-db variation in carrier input is obtained

CENTRAL OFFICE TERMINAL

REMOTE TERMINAL



by the combination of the rectified control signal voltage exceeding the reference voltage of a silicon-aluminum junction diode and by the expansion characteristic of the variolossor. The variolossor uses a specially coded set of the silicon diodes which are matched for both ac and dc characteristics. Modulation products introduced by the variolossor have been kept below those produced in the associated receiving demodulator.

3.7 *Signaling*

The need to transmit customer signaling information over a carrier channel required the development of means of passing dialing and supervision signals toward the central office. It also required passing ringing information to the remote terminal for the types of multiparty ringing generally used in the Bell System, including four-party selective service, eight-party semi-selective service and divided code ringing.

These requirements were met in such a way that the carrier system can be inserted into a normal voice frequency circuit and function without requiring any change in the existing signaling equipment in the central office or in the customer's telephone. The central office terminal is activated by 20-cycle ringing signals which are reproduced at the remote terminal. The remote terminal is activated by switchhook signals and dial pulses which are reproduced at the central office terminals. Thus, the two directions of signaling require completely different circuits. A block schematic of the arrangement used is shown in Fig. 5.

3.7.1 *Ringing*

The customer signaling originating in Bell System central offices consists of 20 cycles superimposed on plus or minus battery and applied between either tip or ring and ground. These signals control the transmission over the P1 carrier system of three in-band frequency tones, the proper combination of two of them serving to select the party to be rung from the far end. The third tone (2,500 cycles), modulated at a 20-cycle rate, carries the information as to whether 20-cycle ringing is present or absent. In-band frequencies were chosen to encode ringing information for transmission over the carrier channel because of the substantially lower cost of in-band filters as compared to those required for out-of-band transmission.

The three signaling tones are generated by three transistor tone oscillators incorporated in a P1 central office terminal. One set of three oscillators can be arranged to supply four central office channel terminals.

The transmission of the tones is controlled by three diode-operated

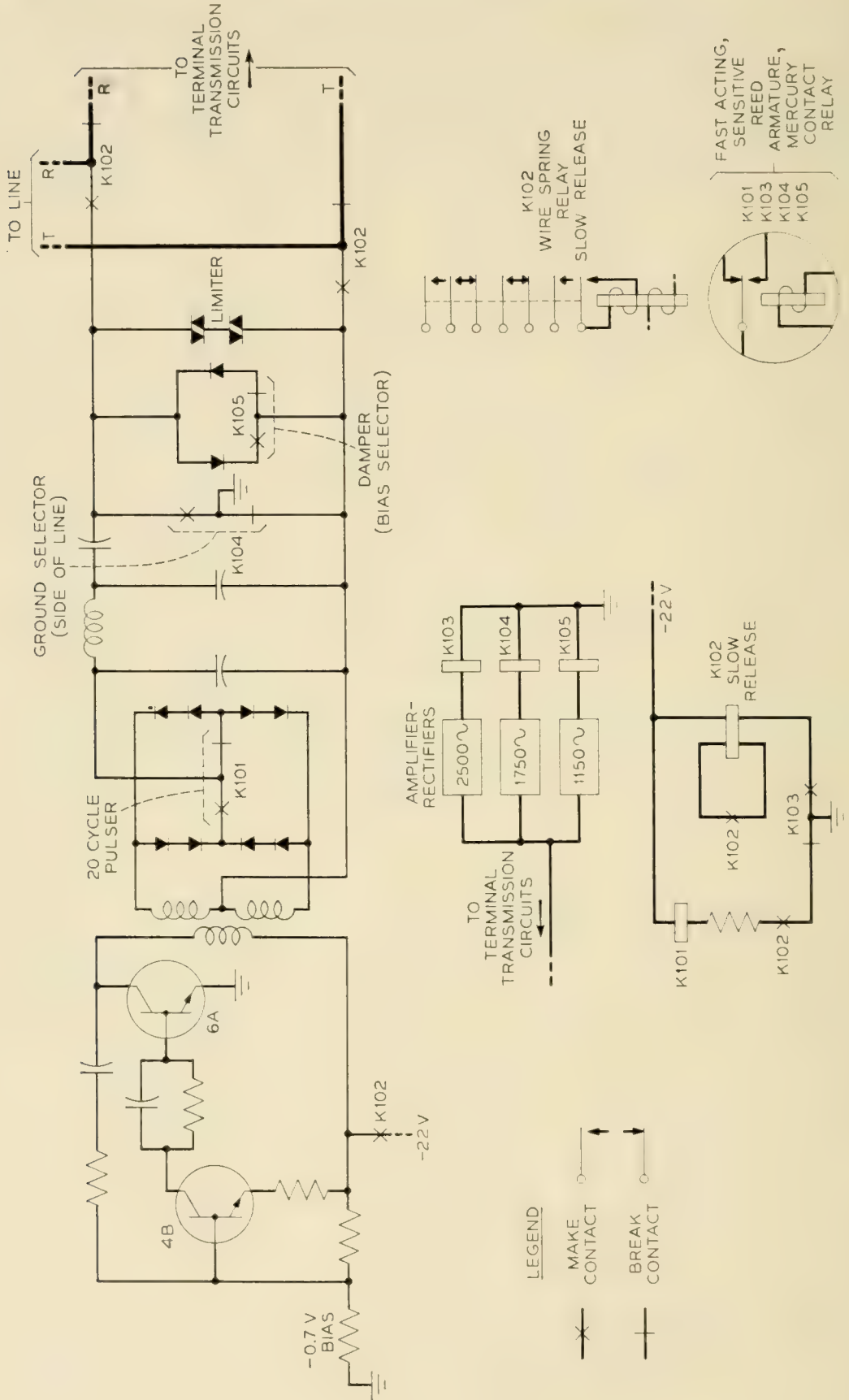


Fig. 6 — Remote terminal 20-cycle ringing generator.

keyers which function independently, depending on the nature of the ringing signal. The 2,500-cycle keyer responds to 20 cycles applied to either the tip or ring conductors. The gating diode is back biased about 3 volts to prevent random noise peaks from operating it. The 1,750-cycle keyer responds to any ringing signal applied to the tip circuit. A diode oppositely poled to the gating diode has a high breakdown so that any reverse voltage peaks leaking through to it will not open the gate. The 1,150-cycle keyer responds to 20-cycle ringing voltage superimposed with either plus bias voltage applied to the tip or with minus bias voltage applied to the ring conductor.

At the remote terminal the signaling tones are each selected at the output of the expander by tuned circuits, amplified by a transistor and rectified to activate relays controlling the customer ringing. The 2,500-cycle tone interrupted at a 20-cycle rate controls the remote ringing generator, the 1,750-cycle tone determines whether ringing is to be on the tip or ring side of the line, and the 1,150-cycle tone whether the bias applied to the line is positive or negative.

The ringing power at the remote terminal is provided by a 3,000-cycle transistor oscillator which uses a 2-watt 6A transistor in the second of its two stages, as shown in Fig. 6. The rectified output of the oscillator is pulsed at a 20-cycle rate by a relay controlled by the 2,500-cycle in-band tone and applied to the line through a low-pass filter to reduce the harmonic content of the ringing signal. Positive or negative bias is provided from one of two clamper diodes. In order to obtain sufficient power from the ringing generator, two electrolytic capacitors are used in a voltage doubler configuration. The ringing generator draws nearly 500 mils of battery current which, by P1 standards, is a heavy power drain. In order to keep this drain to a minimum, the ringing generator is activated only during the ringing period. Also the generator is connected to the customer's line only during the ringing period to remove its shunting effect on the talking circuit.

3.7.2 *Supervision and Dialing*

Carrier on-off signaling was chosen to transmit supervision and dialing signals from the customer to the central office. This method was used because of the ease of implementing it and because of savings in dc power drain at the remote terminal achieved by transmitting the carrier from that terminal to the central office only during the off-hook condition.

An off-hook signal on the voice frequency extension of the remote terminal, caused by the customer lifting his handset, is used at the remote terminal to activate the transmitting amplifier and to remove a short-

circuit from its input. This results in carrier being transmitted to the central office terminal, where it causes relays in the terminal to recreate the customer's line short across the line connecting the central office carrier terminal and the central office switching equipment.

Dialing at the customer's instrument alternately opens and shorts the voice frequency extension at each dial pulse. Opening of the line operates a relay in the remote terminal which short-circuits the transmitting amplifier input and causes the relays in the central office to follow the pulsing of the received carrier, recreate the dial pulses there, and operate the central office switching equipment.

3.8 Repeater

Repeaters are used in the P1 system whenever the line transmission loss exceeds the maximum that the terminals can accommodate. The repeaters use transistors for gain instead of electron tubes, but otherwise are schematically very much like previous repeaters as shown in Fig. 7. The two directions of transmission have similar functions differing only in the frequency band that is amplified and the options that are used in the gain regulating circuits.

Identical high-low pass directional filters at each end of the repeater separate the directions of transmission into the high and low frequency groups. Optional phase correction sections for these directional filters are used along a line to improve the phase characteristics in the cut-apart region. The directional filters are connected to the line through the same line matching coils and secondary protection circuits used in the terminals. Repeater input span pads are used to build out the line loss to its

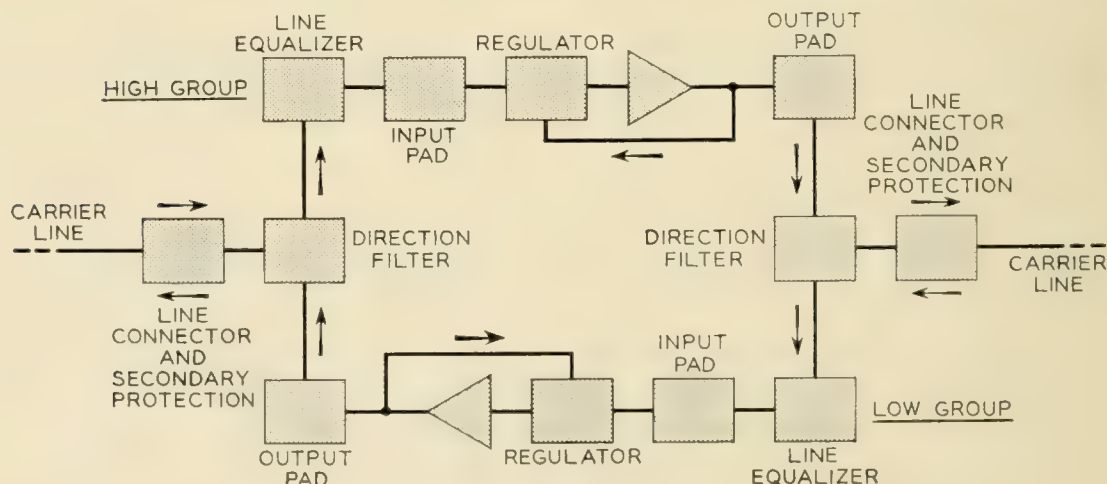


Fig. 7 — P1 repeater block schematic.

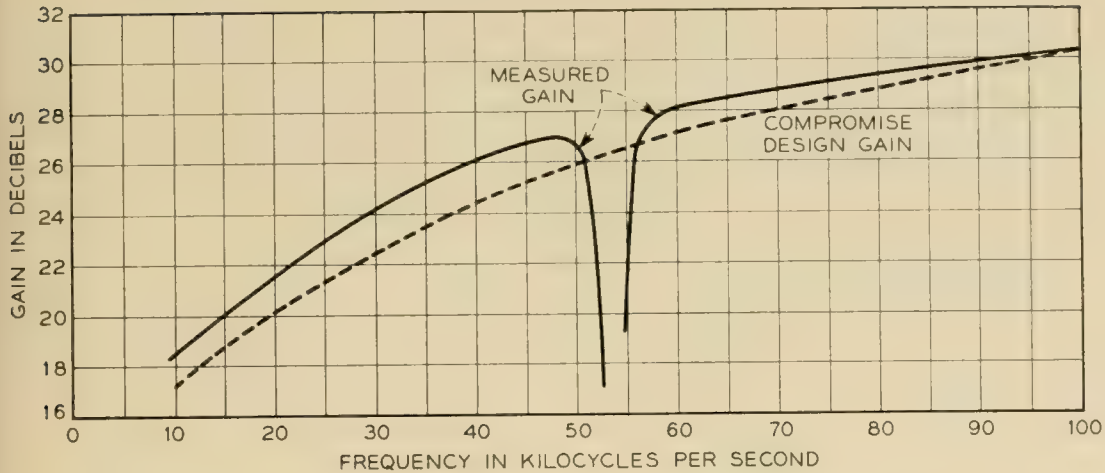


Fig. 8 — P1 repeater gain-frequency characteristic.

proper value for the nonregulated repeater and to adjust the power of the received carriers to the center of the regulator range for a regulated repeater. The output span pads are used where it is necessary to adjust the repeater output power in order to equalize levels between a P1 system and other P1 systems or other types of carrier systems operating on the same open wire line. These pads provide attenuation in 2 db steps up to 30 db.

The repeater amplifiers were designed to have a wide enough frequency band to cover both high and low groups of frequencies, so that each repeater contains two identical amplifiers. Each amplifier has three transistors with each stage connected as a common emitter. Western Electric Company PNP 7B and 6B transistors are used in the first and last stages, respectively, and a NPN type 4C transistor is used in the second stage. Local feedback is required around each transistor to reduce the gain spread and phase variations among units. Overall feedback is obtained around the three transistors with hybrid coils at input and output.

The repeater equalizer characteristic represents a compromise for several types of transmission facilities generally encountered in the rural plant. The equalizer design also covers both the high and low frequency groups so that identical equalizers are used for both the high group and low group sides of the repeater.

A preliminary characteristic for the overall repeater gain is shown in Fig. 8 plotted against the design objective. There is a significant departure in shape only at the cut-apart frequencies. This will be corrected sufficiently to permit as many as four repeaters to be used in tandem. As the design objective is a compromise of the loss of several types of lines that may be encountered in the use of this system, the departures

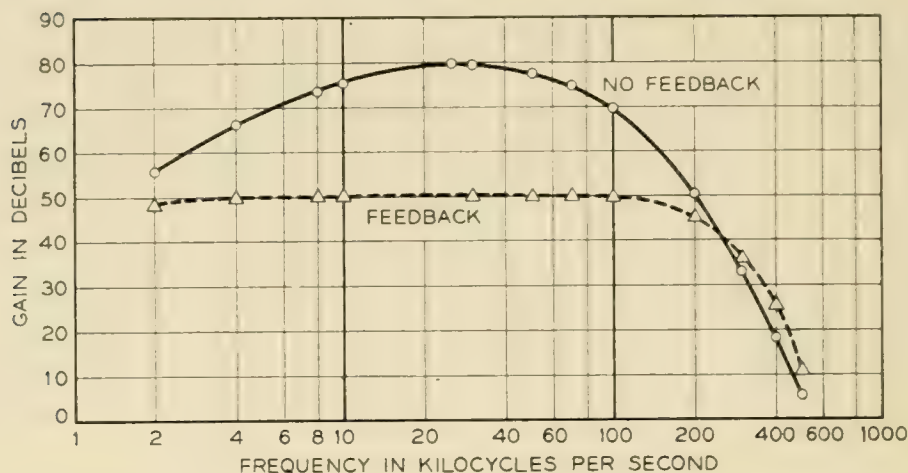


Fig. 9 — P1 Repeater amplifier gain-frequency characteristic.

in system performance may vary considerably from that shown. The repeater amplifier gain frequency characteristic, plotted in Fig. 9, shows a non-regenerative peak gain of the three-stage transistor amplifier of 80 db and a feedback gain characteristic of 50 db. This provides 20 db or more of feedback over the transmitted band to produce the necessary operating stability with temperature and power supply variations, and a working value of modulation suppression.

3.9 Repeater Regulation

Repeater regulation will be furnished as an option where variations in line loss exceed the terminal regulating range. It will usually be necessary on systems employing more than one repeater in order to control noise performance. Repeater regulation in the direction of transmission from central office to remote terminal is controlled by the total carrier power of the channels working on one system. In the opposite direction, the repeaters will regulate on a low level carrier frequency pilot because the channel carriers are not always present in that direction of transmission due to their signaling function. The pilot frequencies are shown in Fig. 1.

The repeater regulator, shown in schematic form in Fig. 10, functions in much the same manner as the terminal regulator. The principal differences between the two regulators arise from the requirement that interchannel modulation must be appreciably less than 1 per cent in the repeater. To limit the contribution of the repeater regulator to a small value, the variolossor operates into a lower impedance and at a higher control current than used in the channel regulator.

The input section to the control amplifier is either a flat bridging pad for the case where all carriers are always present on the line or a pilot pick-off filter and its associated single transistor amplifier where carriers are turned on and off for supervision. The latter extra amplifier is necessary because the pilot power is 20 db below the power in each normal carrier.

The regulator stiffness provided by the repeater regulator results in a variation of 1 db in output carrier power for a 10-db variation in received total carrier or pilot power.

4. COMPONENTS

Development of the passive components of the P1 carrier system, including the various filters and other networks, were influenced by three major considerations. The manufacturing cost had to be as low as possible consistent with the traditional standards of Bell System service life. The components had to lend themselves to maximum utilization of the printed wiring techniques to be used as the basic equipment method. And lastly, advantage wherever possible was to be taken of the fact that transistors are low power devices.

4.1 Filters

Component-wise, filters are the most important single assembly determining the first cost of a carrier system employing frequency division multiplexing and frequency separation for obtaining equivalent four-wire

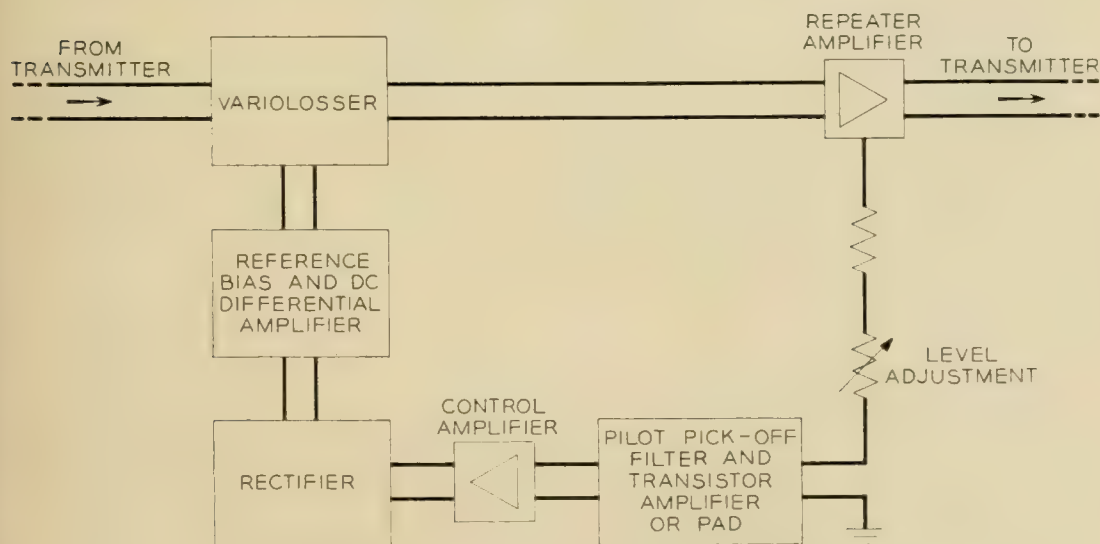


Fig. 10 — P1 Repeater regulator block schematic.

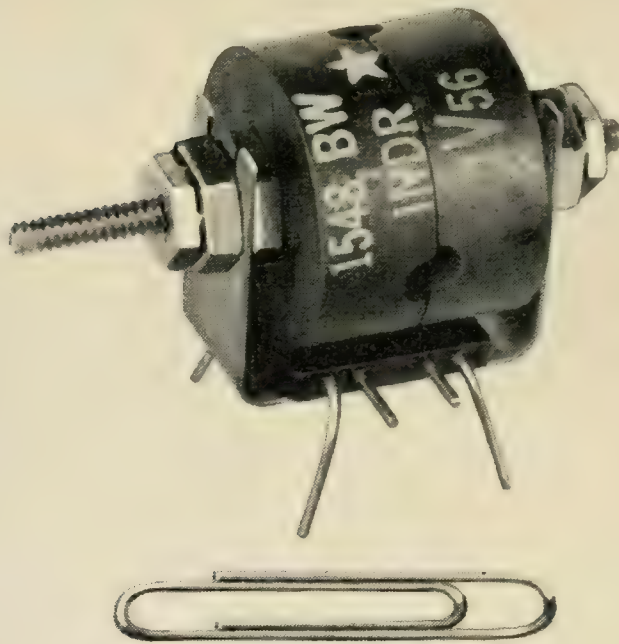


Fig. 11 — Miniaturized inductor for P1 carrier.

operation on the line. Past experience has shown that one-quarter or more of the first cost is often chargeable to the various filter. The decision to employ double sideband modulation was largely based on the knowledge that when frequency space is available the double sideband channel filters are generally the least expensive.

With this decision made, every effort was directed toward the achievement of channel filter designs of maximum efficiency in element utilization. Inexpensive wide-limit capacitors were used, and the desired performance achieved through the use of an adjustable ferrite inductor expressly developed for the P1 system. The filters are rapidly adjusted in the manufacturing process using visual display testing circuits.

4.2 Inductors

The inductor which makes this possible is shown in Fig. 11. It is designed for printed wiring use and provides a wide range of inductance while maintaining excellent "Q" performance in the carrier and voice range. This is accomplished in a single basic design by so selecting the winding for particular nominal inductances that the air-gap adjustment remains at or near its most efficient setting. Inductors of this type were

used not only in all the channel, demodulator, and signaling filters in the terminal and in the directional filters at the repeater, but also in the channel oscillators and other parts of the circuitry where an inexpensive, adjustable element offered manufacturing or service advantages.

4.3 Capacitors

Most of the wide-limit capacitors used in the filters are of the commercially available molded mica type. Where the capacitance values would require large and expensive mica units both in filters and other parts of the circuit, newly available foil-Mylar capacitors were used. These take the form of very small pigtail units in a range of physical sizes similar to those of the solid tantalum capacitors described below. The Mylar capacitors have low working voltages in these miniature sizes and can



Fig. 12 — Prototype solid tantalum capacitors for P1 carrier.

be employed because of the low voltage protection provided for the transistorized circuits. Both cost and space savings are realized in these capacitors since no cans or potting are required due to the stability of the Mylar dielectric under moisture exposure.

Another new type of capacitor has found widespread use in the P1 system. This is a solid tantalum electrolytic capacitor used in place of the usual paste or liquid electrolytic capacitor. The solid electrolyte is manganese dioxide deposited upon the capacitor surfaces. The anode is made from tantalum metal and upon its surfaces is deposited the tantalum oxide which forms the dielectric. The cathode is an enveloping metal completing the capacitor structure. This new design of capacitor is now available in values up to 100 microfarads in a very small volume. It is expected to be less expensive than other electrolytic capacitors while at the same time providing a rugged structure which is relatively inert electrochemically and which has better stability in operation and storage. Fig. 12 shows prototype models of typical solid tantalum units.

4.4 Transformers

Transformer needs in the P1 system are met by two miniature structures which were made possible by the use of low power transistor circuits. The carrier frequency units employ a manganese zinc ferrite core, a spool winding and wire terminals which permit assembly on printed

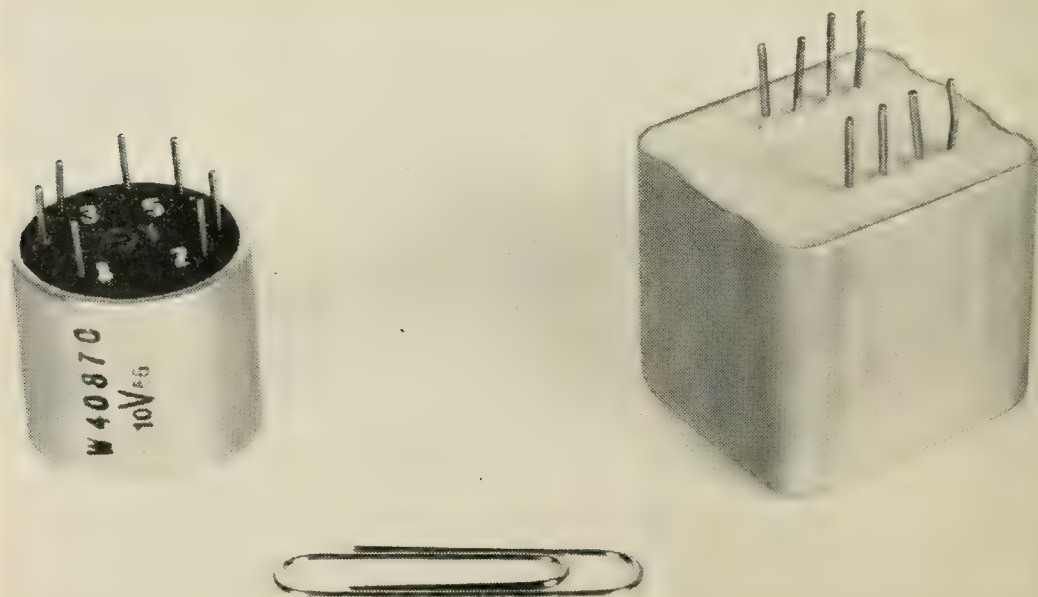


Fig. 13 — Carrier and voice frequency transformers for P1 carrier.

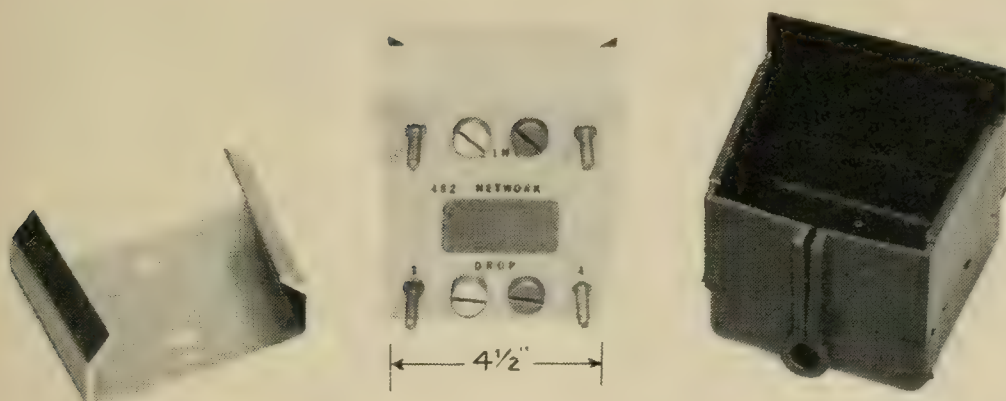


Fig. 14 — Example of P1 carrier line network.

wiring boards. They are potted with an asphalt compound in a cylindrical aluminum can. The voice frequency transformers are wound on laminated core structures of permalloy. The units are potted in an epoxy resin in rectangular aluminum cans. The terminal plate carrying the wire terminals for mounting is a cast unit of a styrene polyester. Both types of transformer are shown in Fig. 13.

4.5 Line Networks and Filters

Also deserving mention is a new series of line networks and filters (which do not form part of either the terminal or repeater equipment) with specific functions described in Section 7. All of the networks have been designed with the same type mounting arrangement shown in Fig. 14 with two sizes used depending on the number of components housed. The networks are cast in a styrene polyester. High voltage protection is self-contained and sturdy terminals are provided for bridle wire connection. By means of side slots in the casting the network is mounted on a wedge-shaped holder which is fastened to the crossarm or pole. A flexible rubber cover is snapped over the face of the network to protect against weather effects.

5. EQUIPMENT ARRANGEMENTS

The emphasis placed on economy in this development project made it necessary to consider a number of different approaches before deciding on the physical arrangement provided for both central office and pole mounted equipment. At both locations the terminals for each channel

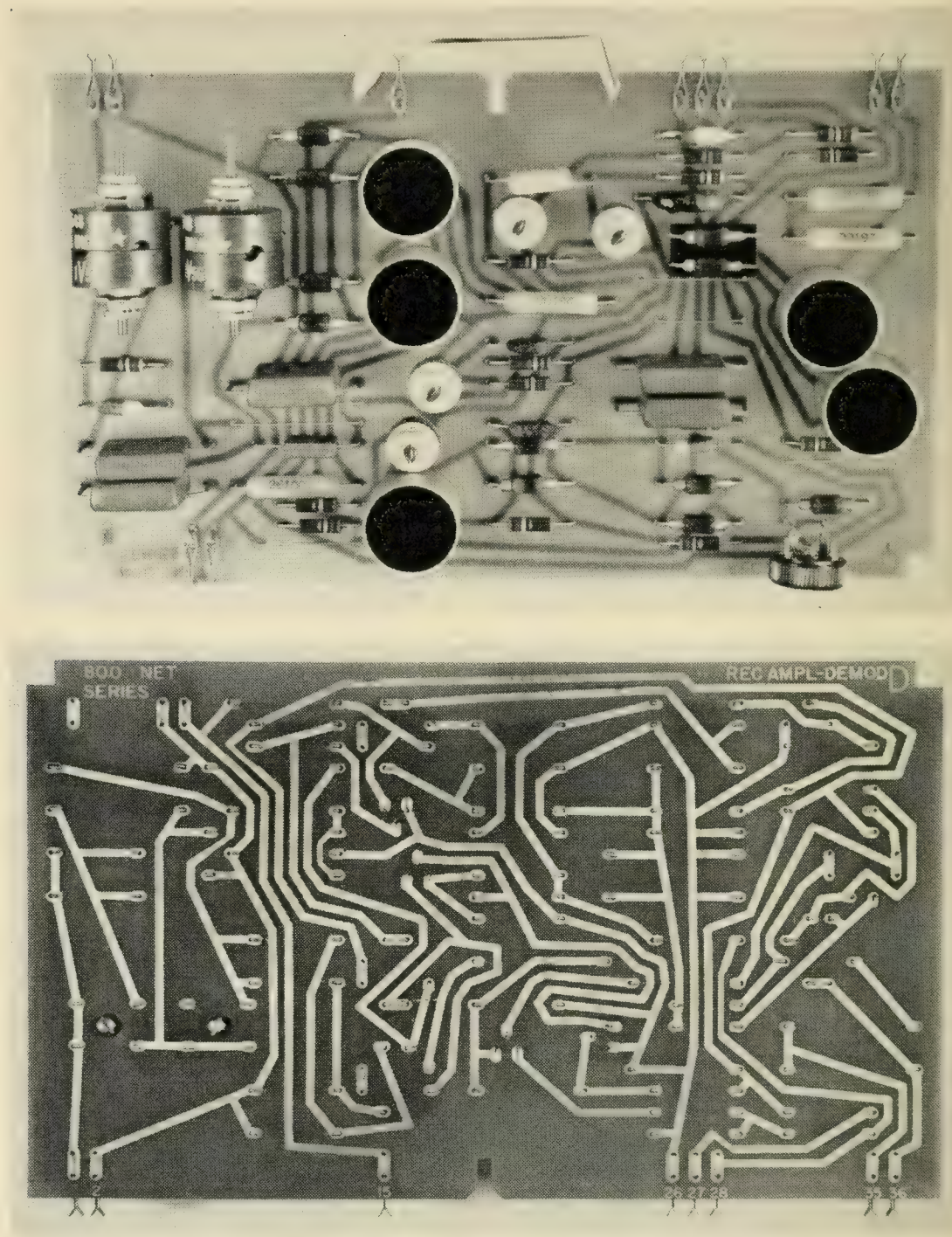


Fig. 15 — Front and back of typical printed wiring board.

were treated on an independent basis, thus providing maximum flexibility in application. While the use of transistors made it possible to take advantage of miniaturized components, the major emphasis in design has not been on miniaturization; instead it has been to achieve low manufacturing costs, simplicity of engineering and installation, and a minimum of maintenance effort. The recent trend toward automation in

manufacture of electronic equipment has also influenced the design to a great extent.

5.1 *Printed Wiring Boards*

To best meet these objectives, use has been made of plug-in units which have proved successful in other carrier systems, such as the N1 and O. The assembly technique used here, however, is an entirely new approach for carrier equipment in that the plug-in unit consists of a printed wiring board on which all components are mounted. Printed wiring, which is a comparatively new engineering technique, was selected because of its applicability to automatic assembly, including mass soldering of connections. In addition, the use of printed wiring greatly simplifies testing and inspection and assures a more uniform product. The two sides of a typical printed wiring board are shown in Fig. 15.

5.2 *Interconnection of Boards*

The interconnection of the various plug-in units or printed wiring boards, required to make up a complete P1 terminal or repeater, is accomplished by means of a wire connector specifically developed for this project. Basically, the connector consists of a number of accurately spaced bare wires running parallel to each other and imbedded in cross member strips of insulating plastic material. At fixed intervals the wires are exposed, and this is where contact is made to terminal connectors mounted on the printed wiring boards. These terminal connectors, shown in Fig. 16, are made of spring tempered phosphor bronze

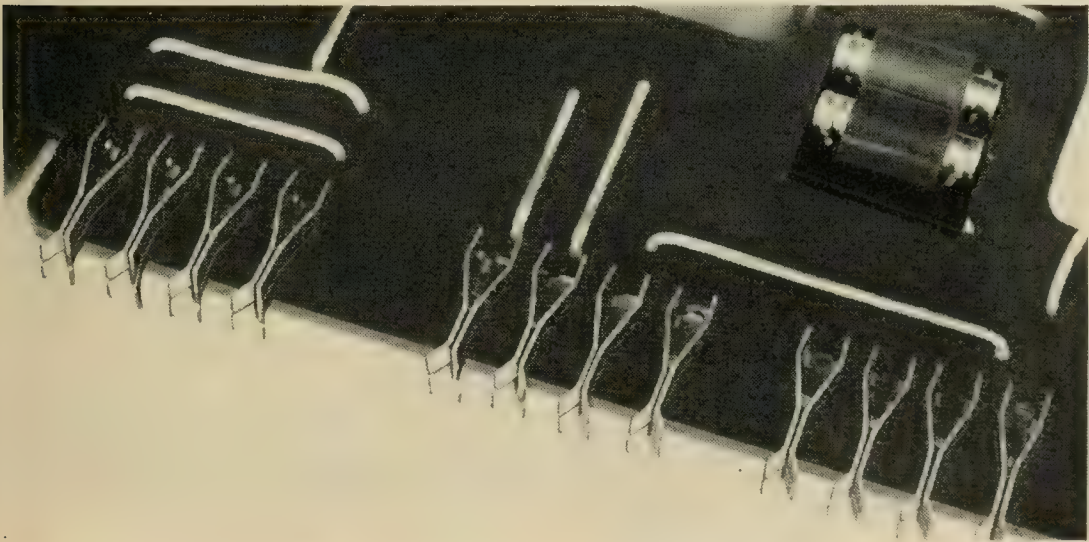


Fig. 16. — Closeup of terminal connectors.

and consist of bifurcated cantilever springs, providing a total of four contacts for each connection.

As can be seen in Fig. 17, the wire connector is actually a molded phenolic box into which are inserted all of the printed wiring boards that make up a complete terminal or repeater. The terminal connectors on the boards thus engage the wires that are imbedded in the back of the connector. To insure contact reliability, a finish of precious metal is provided on both the wires of the connector and the terminal connectors on the board. Additional flexibility in the interconnection of the boards is

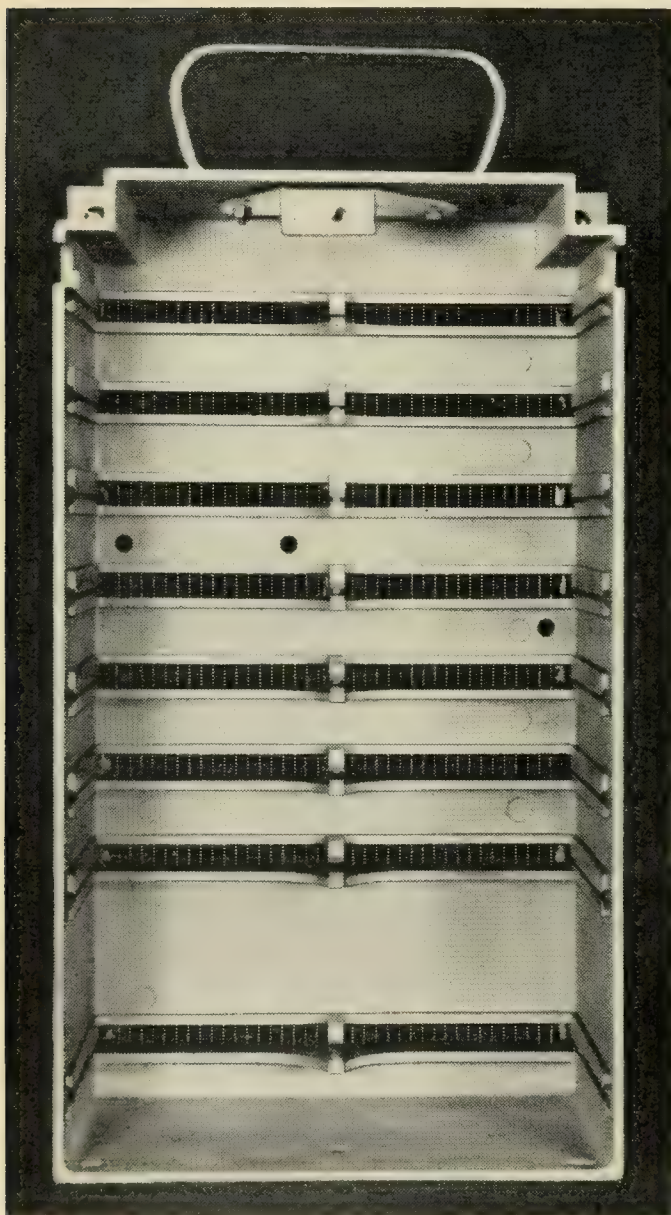


Fig. 17 — Prototype model of connector box unequipped showing grid wires.



Fig. 18 — Typical terminal network mounted in a prototype connector box.

obtained by cutting the wires at various points by simply drilling holes in the phenol structure supporting the wire.

5.3 *Terminal and Repeater Mounting*

A complete terminal ready for installation at a remote location is shown in Fig. 18. The top position in the connector is shown vacant. This is where the connections to line and power supply are made by means of another plug-in printed wiring board with attached flexible wiring for the external connections.

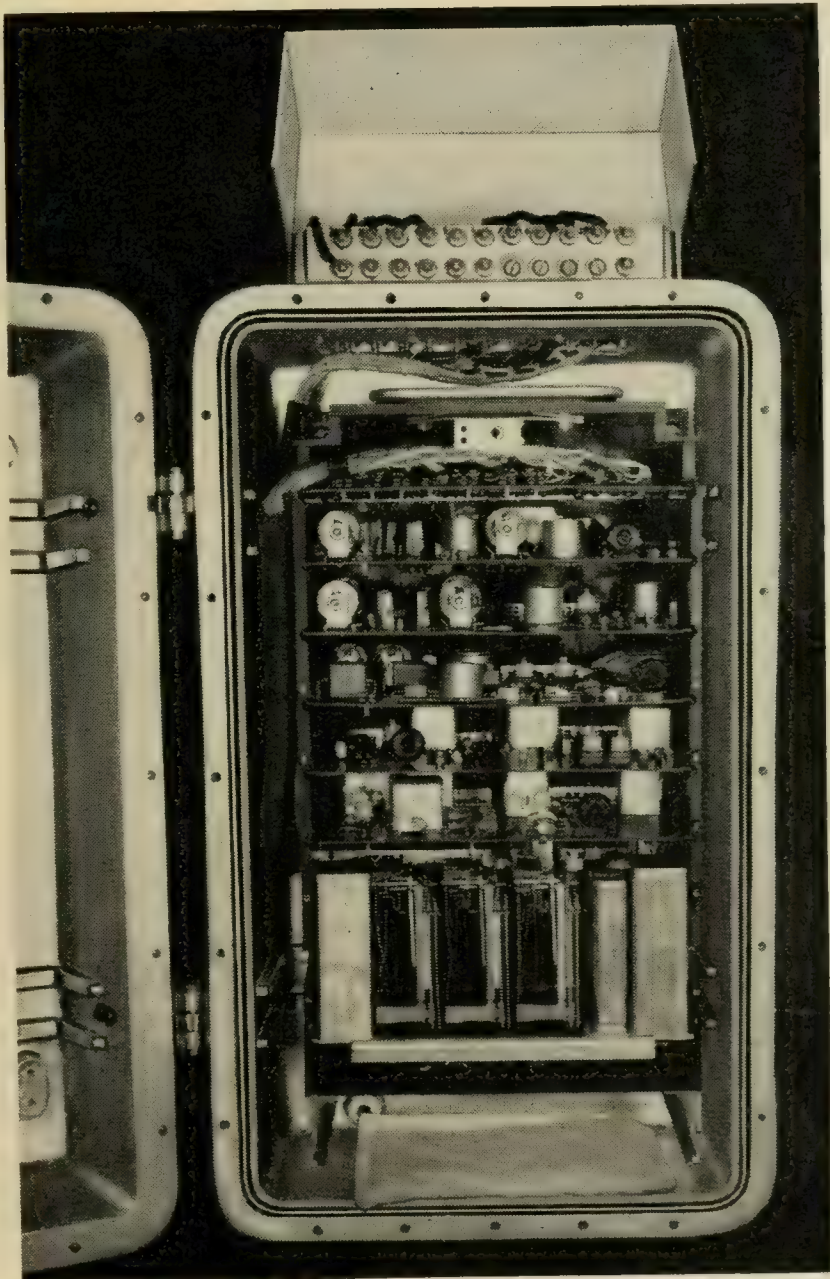


Fig. 19 — P1 carrier remote terminal in pole mounted cabinet.

The equipment described here is equally adaptable to central office mounting and pole mounting at remote locations. At the remote locations, however, it is necessary to provide the equipment with an outer housing which gives protection from all kinds of weather and even from moisture condensation. The opened housing is shown in Fig. 19. Fig. 20 shows a typical remote mounting of the housing on the left. In previous electron tube carrier systems the amount of heat generated by the equipment itself was sufficient to prevent moisture condensation. In the case

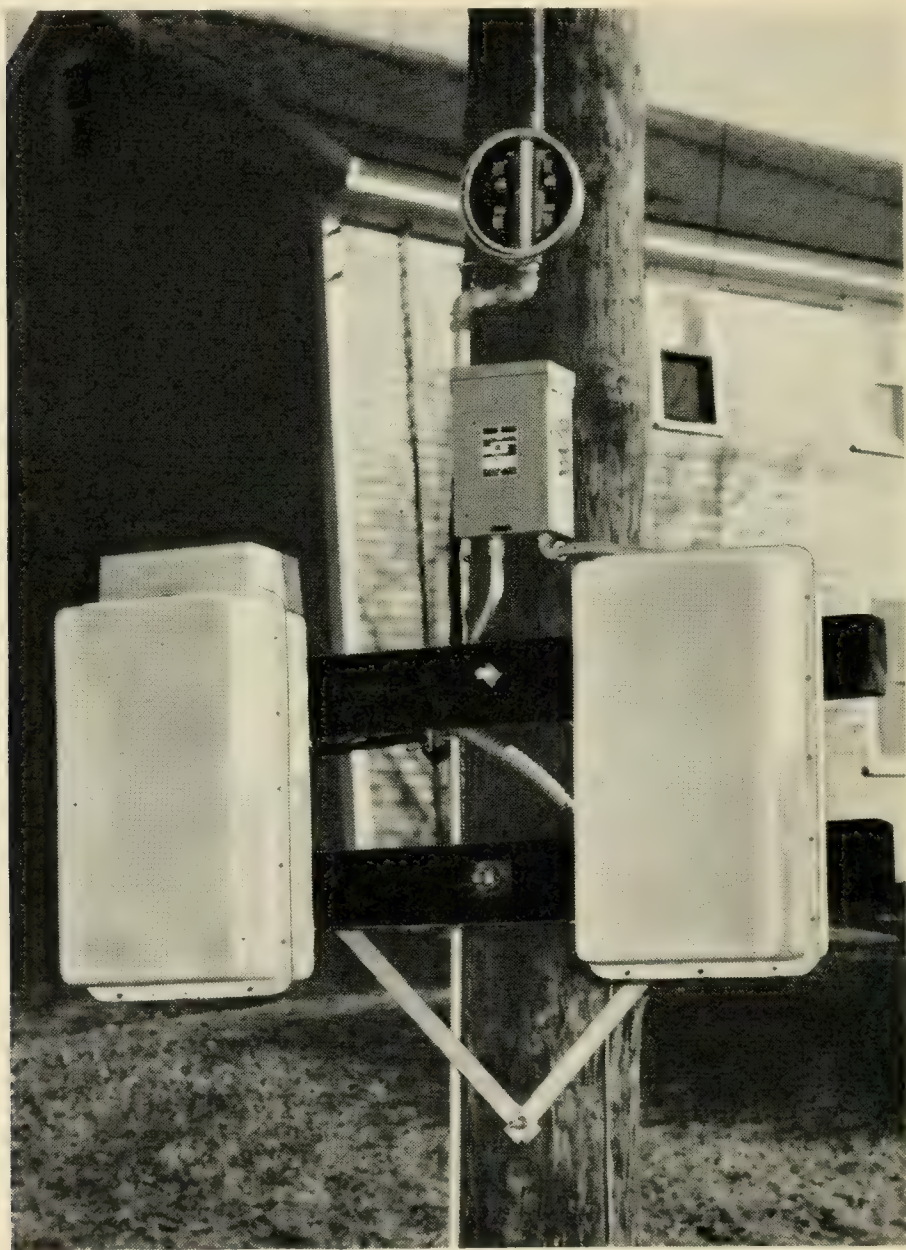


Fig. 20 — Example of remote mounting of P1 terminal and ac power supply.

of the P1 carrier, however, the heat dissipation during the idle period is less than 1 watt for the entire terminal. To prevent condensation, the housing or apparatus case is sealed by means of a neoprene gasket. To further reduce the moisture content of the trapped air, the use of a desiccant is specified. The apparatus case is made of die-cast aluminum with the outside walls finished in white enamel to keep heat absorption to a minimum. The system was designed to operate between temperature extremes of -40°F and $+140^{\circ}\text{F}$. This limitation might necessitate the additional installation of sun shields in a few cases where extreme temperatures prevail.

The terminal equipment at the central office makes use of the same type of printed wiring boards plugged into a connector as used at the remote location. In the central office, however, the outer housing is dispensed with and the connector is mounted on mounting brackets on standard relay racks. The relay rack layouts can be arranged in a number of ways to suit the particular installation, since no shop wired bays are used. A typical 11'6" relay rack layout will provide for 10 terminals. No line jacks or alarm features are provided and fusing may be obtained from existing fuse boards in the office. The equipment also lends itself to wall mounting in locations where relay rack space is not available.

5.4 Testing and Maintenance Features

One great advantage of the equipment design used in the P1 carrier system is the ease with which an entire terminal or repeater can be transported to, and installed at, a remote location. In case of trouble, the entire equipment unit, be it a terminal or a repeater, can be readily replaced. It is not expected that the maintenance man will attempt to replace an individual printed board at a remote location; however, this procedure is perfectly feasible in a central office. To facilitate the location of trouble in a unit, the various boards are provided with test points located at the outer end of the boards so as to be easily accessible to the maintenance man.

Certain precautions will have to be taken at central repair centers in replacing defective individual components in order not to damage the printed wiring. Too much heat applied by a large soldering iron will destroy the adhesive bond between the copper conductor and the phenolic board, but repair can be made under certain controlled conditions. A limited amount of wiring modifications can also be made to the printed wiring by inserting strap wires in place of components.

6. POWER SUPPLIES

The design of a carrier system with low power drain made possible the development of a low-cost, reliable dc power supply for the carrier equipment. Because the central office carrier terminal was designed to utilize standard central office voltages (24 or 48 volts), only the power supply for the remote equipment will be described here.

Early exploratory studies showed that conventional power supply designs would miss the first and annual cost objective by an uncomfortable margin. A number of unconventional approaches were studied:

- (a) Storage batteries charged over the carrier line.
- (b) Storage batteries placed in service with full charge and removed to a central point for recharging.
- (c) Solar power plants.
- (d) Wind power plants.
- (e) Thermoelectric power plants.
- (f) Dry cells.

In all of the above cases the power plant was either too costly, too large, or technically unfeasible, and none could prove in over the conventional conversion of ac to dc where commercial power is available. This was true despite need for a storage battery to operate the system during ac power failure intervals and to provide peak ringing power.

6.1 AC Rectifier-Storage Battery Plant

The basic elements of the power plant circuit, as shown in Fig. 21, are the conversion section represented by the step-down transformer T1 and

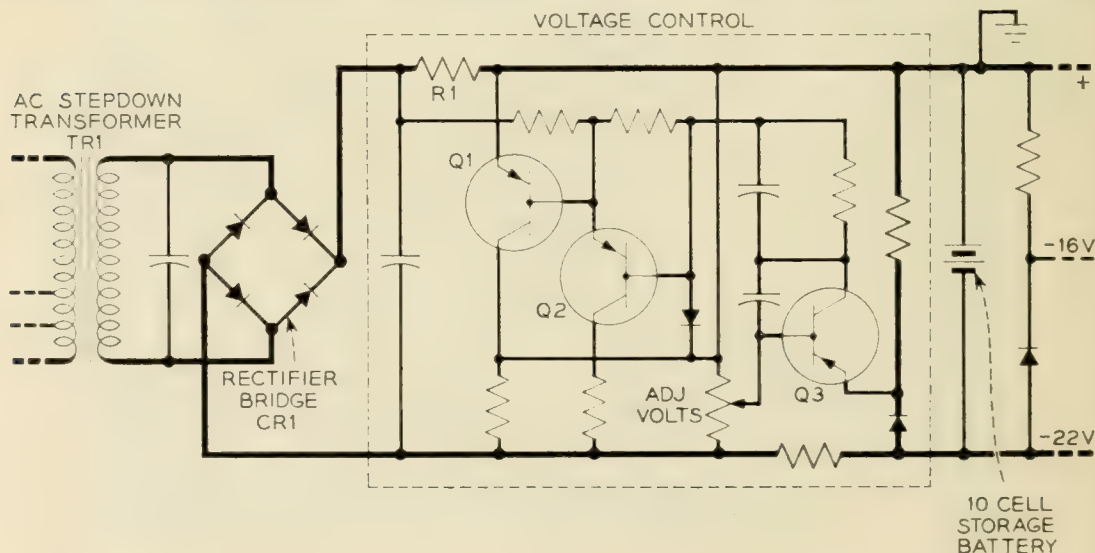


Fig. 21 — Schematic of ac rectifier-battery power supply.

the semiconductor rectifier bridge CR1, the voltage control circuit represented by that part of Fig. 21 enclosed in dashed lines, and the energy storage circuit represented by the battery.

Rectification is obtained with germanium rectifiers that are very efficient, have long life with negligible aging, and are very compact physically. The output of this rectifier is not constant, because the output voltage will vary with the ac input voltage and the dc load current drawn by the carrier terminal. Thus a regulating circuit must be provided.

The regulating network senses the voltage across the battery and compares this voltage to a reference obtained from a silicon junction diode biased in the reverse direction.³ Any error in the output voltage is converted to a current signal in the first amplifier stage and amplified by the second stage transistor Q2. The amplified error current is then used to control the impedance of transistor Q1 which acts as a current shunt around the battery.

The fundamentals of the operation of this regulating system are shown in Fig. 22. If the load voltage is too high, the network adjusts the resistance of transistor Q1 so that some of the rectifier output current is shunted around the load. The load voltage will then return very quickly to the regulated value. Because the rectifier circuit must not be overloaded by a discharged battery, some form of current limiting must be provided; this is automatically taken care of by resistor R1. The rectifier is capable of supplying indefinitely the current that would be drawn to charge a battery after a very long power failure.

The storage battery is shown in Fig. 23 near the bottom of the power plant housing. It is a new design with a high specific gravity sulphuric

³ D. H. Smith, Silicon Alloy Junction Diode as a Reference Standard, A.I.E.E., Communication and Electronics, No. 16, pp. 645-651, Jan. 1955.

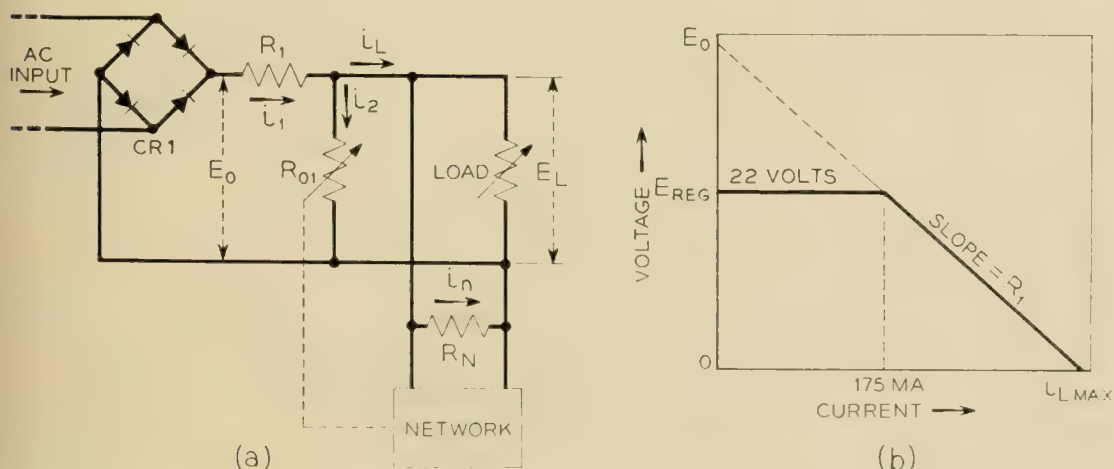


Fig. 22 — Simplified schematic and regulation characteristic of ac power supply.

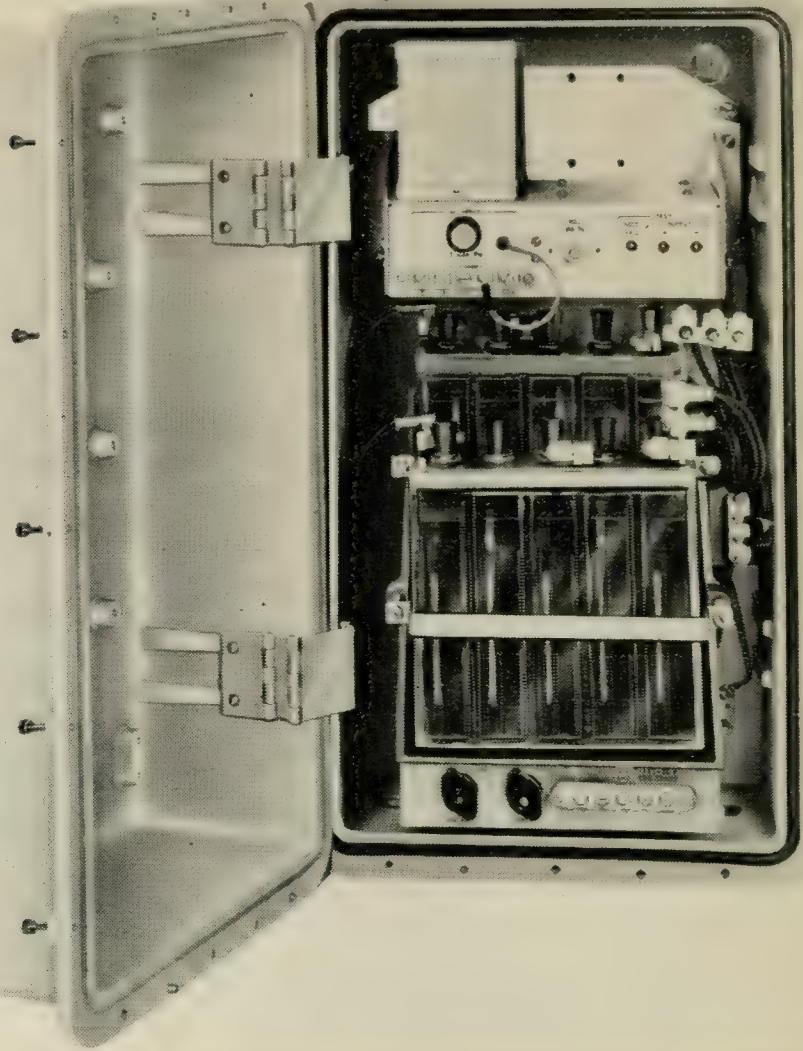


Fig. 23 — P1 carrier ac power plant in cabinet for pole mounting.

acid electrolyte for good low temperature operation and lead calcium alloy electrodes for long life. The battery has 10 cells housed in two jars of five cells each. It is operated at about 23.5 volts and weighs about ten pounds. It provides about six days' reserve for a remote terminal or about two days for a remote repeater. The battery should not freeze at temperatures as low as -40°F , but the storage capacity may be reduced 90 per cent at this extreme. The battery is mounted on steps so that the electrolyte level can be seen through the transparent plastic battery jars. Fig. 23 also shows the compact packaging of the entire power supply within the same type of aluminum housing as used for the carrier terminal. A typical pole mounted installation is shown in Fig. 20. Fig. 24 is a

close-up of the bottom of the rectifier chassis which shows the regulating network mounted on a printed wiring board.

6.2 Air Cell Primary Battery Plant

Because ac will not be readily available at all remote locations, an alternate power supply has been developed and this is shown in Fig. 25.

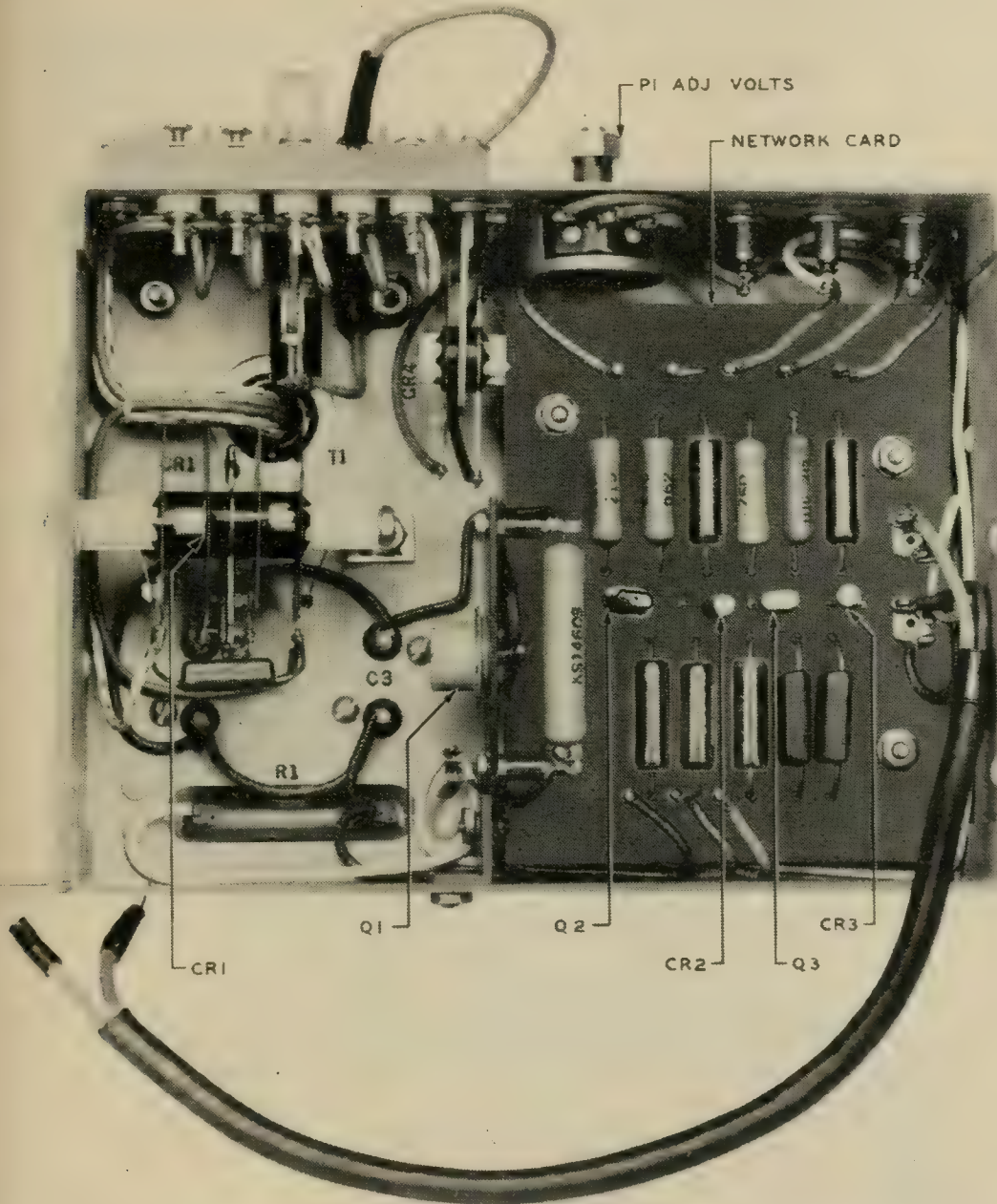


Fig. 24 — Close-up of bottom of rectifier chassis.

The alternative supply uses oxygen-depolarized primary cells having an alkaline electrolyte, and has been used for many years in railway signaling circuits and in the telephone plant. Sixteen battery cells are connected in series to provide enough power for three years of operation of a remote terminal or about one year for a remote repeater. The battery is discarded when fully discharged and is then replaced by a new battery.

7. APPLICATION OF P1 CARRIER TO RURAL TELEPHONE LINES

The P1 carrier system is to be applied to normal exchange loop plant facilities engineered in accordance with the present Resistance Design

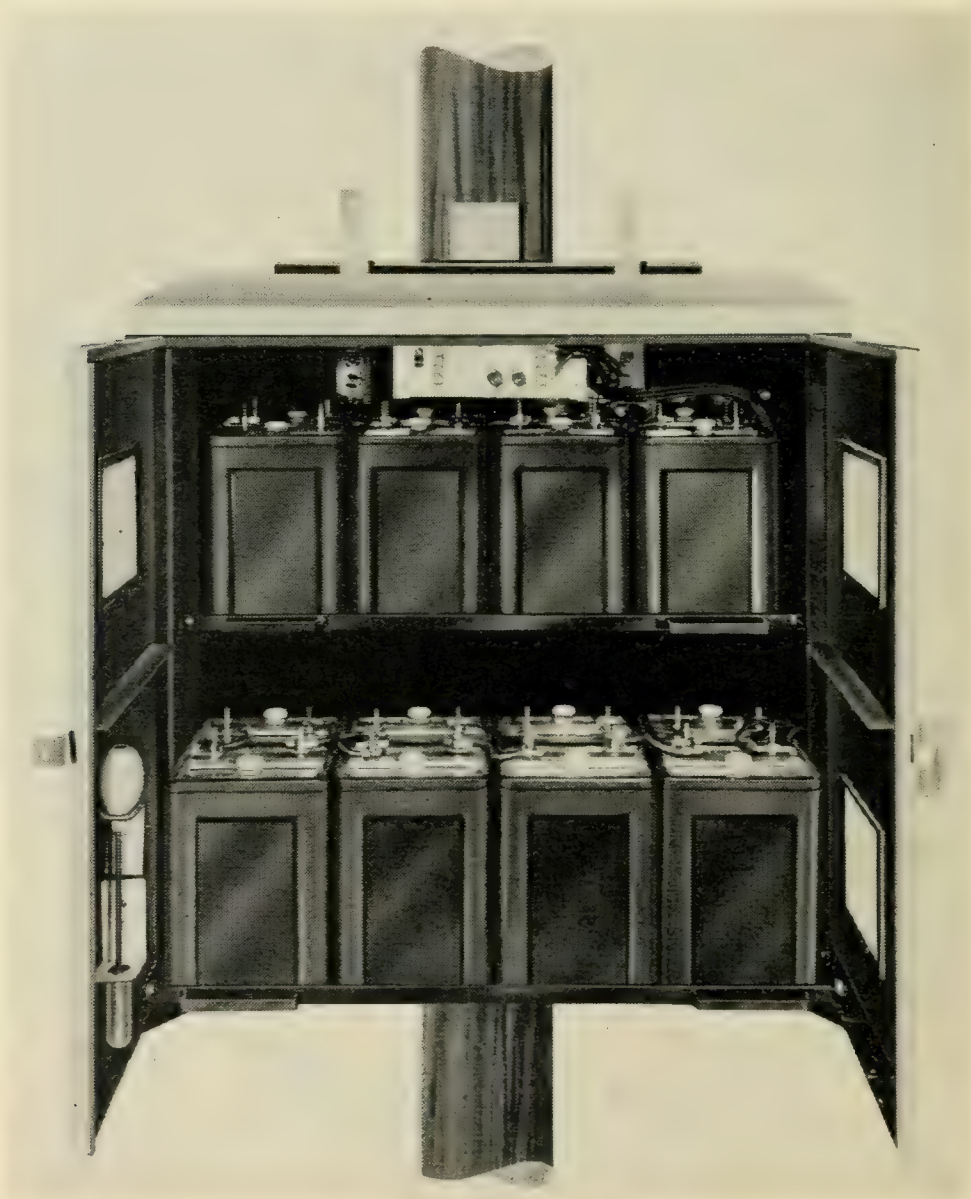


Fig. 25 — Primary battery power plant.

Methods used generally throughout the Bell System.⁴ These facilities consist of mixed gauges of high capacitance cable extended at their outer ends by 109-mil steel, 104 mil-copper or copper steel open wire.

In engineering the carrier line design or carrier layout, the Plant Engineer will determine the carrier line layout necessary to meet the over-all requirements for a suitable carrier transmission path on the available physical facilities. To do so he need only be familiar with the general capabilities of the carrier system, its basic "building blocks," and the limitations that must be considered in applying the system to the physical line. The capabilities of the carrier system have been described in earlier sections. From those descriptions it can be seen that the basic "building blocks" for a P1 carrier system are:

1. Central office channel terminals
2. Remote channel terminals
3. Repeaters
4. Ac or dc remote terminal and repeater power supplies
5. Carrier line networks and filters

A carrier application of these "building blocks" is shown in schematic form in Fig. 26.

The low-pass filters or carrier blocking networks shown are placed at the junctions of the carrier line and side leads of customer drops served by physical or derived voice frequency circuits on the base carrier facility. These filters are required to reduce the bridging loss of the side leads at carrier frequencies and to keep carrier frequencies out of the customer drops to prevent annoyance to the customers. High-pass filters are provided to make the carrier line continuous at carrier frequencies, but divide it into isolated sections for voice frequency distribution.

In addition to these blocks, an autotransformer may be required at the junction of the open wire and cable. The autotransformer, either alone or in conjunction with a junction line filter, is required to eliminate reflection losses and reduce crosstalk at carrier frequencies due to impedance mismatch between the cable and open wire. The junction line filter is required to allow the carrier and physical voice frequency circuit to be used on different pairs in the cable and on the same open-wire pair beyond the cable-open-wire junction. This is necessary where the physical circuit is so long that load coils are required on the voice frequency cable pair and non-loaded cable pairs are required for carrier. A pair of junction line filters may also be used to provide a voice frequency by-pass around a repeater. As illustrated in Fig. 26, this may be necessary

⁴ L. B. Bogan and K. D. Young, *Simplified Transmission Engineering in Exchange Cable Plant Design*, A.I.E.E. Communication and Electronics, No. 15 page 498, Nov. 1954.

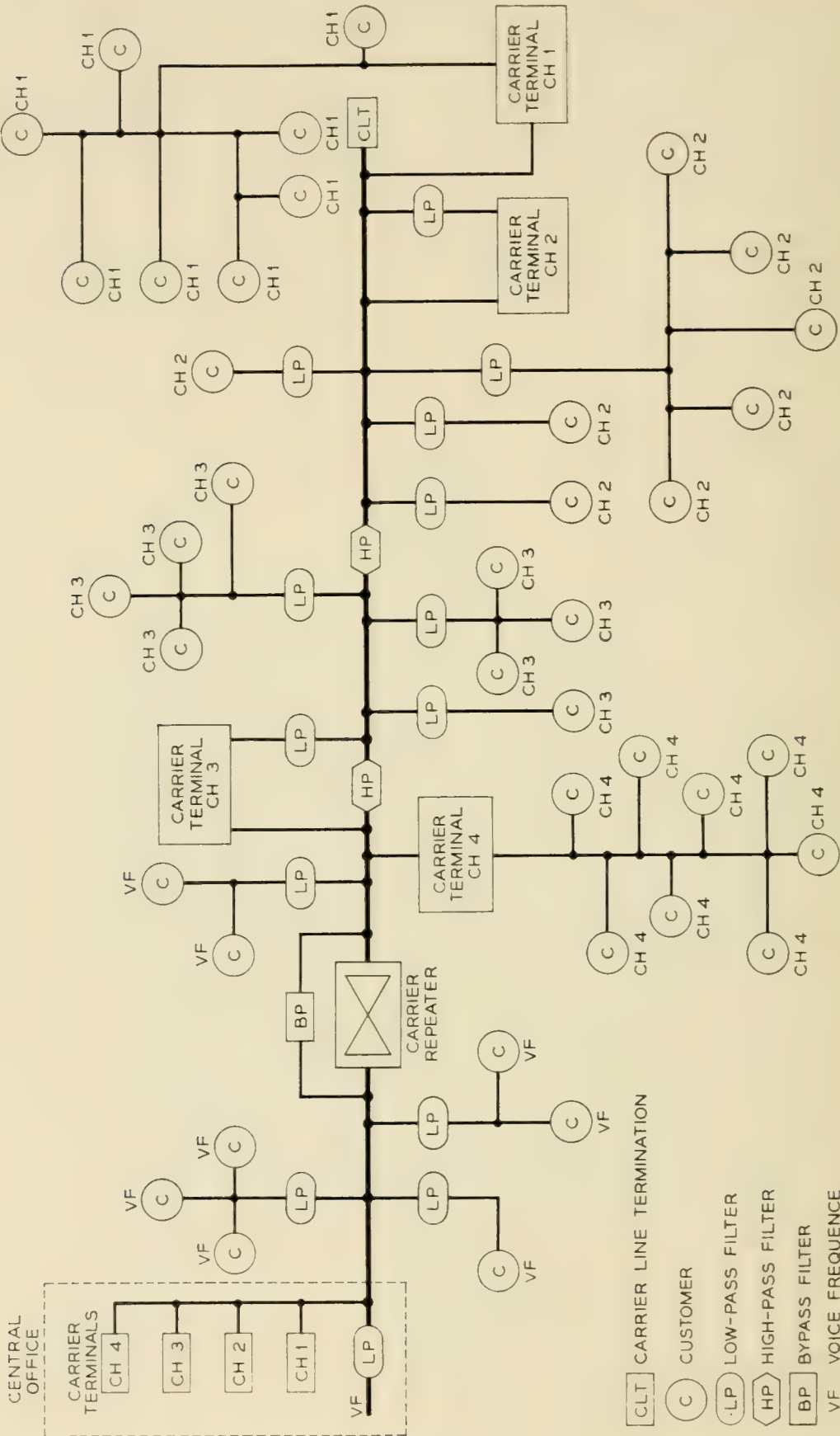


Fig. 26 — Typical P1 carrier application schematic.

to serve customers beyond that point from physical voice frequency circuit.

A carrier line termination network is also provided to terminate the end of the carrier line at all frequencies and thus prevent reflections from interfering with the transmission at remote carrier terminals spaced along the line. This network and all of the other line networks are available in the pole or crossarm mounted arrangement shown in Figure 14 and described in Section 4.5.

Fig. 26 also gives examples of two types of subscriber distribution beyond the remote carrier terminals. One, wire distribution, is indicated by the voice frequency extensions of Channels 1 and 4 and the other, filter distribution, is shown for Channels 2 and 3. Filter distribution permits the carrier line to be used simultaneously for carrier transmission and voice frequency distribution of the derived voice frequency circuit, thus saving the pair of wires required if wire distribution were used.

7.1 *Layout Procedure and Ground Rules*

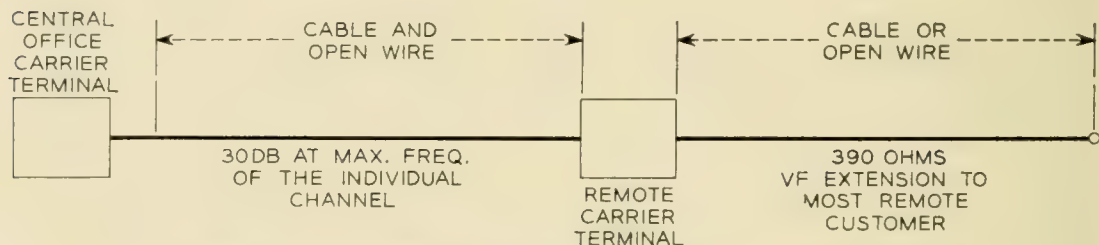
P1 carrier channel layouts for a given rural line will be based on the forecast of commercial requirements for that route. The Plant Engineer must determine the number and arrangements of channels which can be applied within the system limits to meet that forecast. The locations of remote terminals are then chosen based on customer locations, channel frequency arrangements, and the availability of commercial ac power. With the terminal locations fixed, the line losses are determined at appropriate frequencies and repeaters are specified as necessary along with any line networks and filters required for the layout.

The characteristics and limitations of the P1 system lead to certain simplified ground rules which may be used in laying out the carrier channels. Some of these rules are summarized in Fig. 27. The stackable frequency arrangement is used for non-repeatered operation, and the design of the carrier channels permits the bare line loss of each individual channel to be 30 db at the top frequency between the central office terminal and the remote terminal.

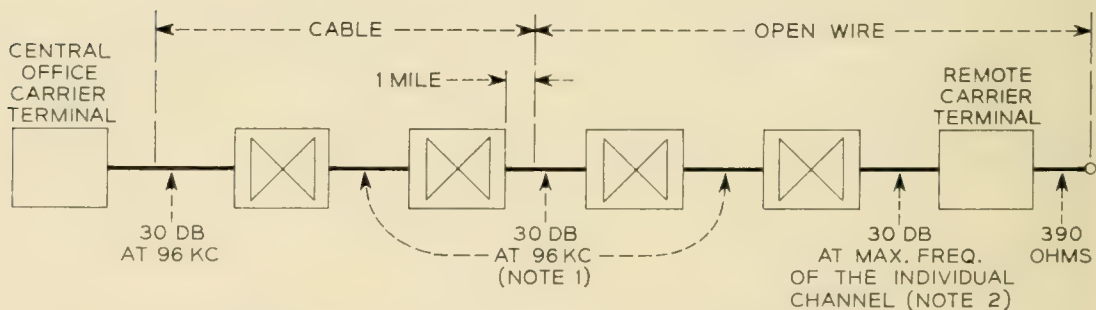
Another limit shown in the figure is that the dc loop resistance of the voice frequency extension beyond the remote terminal can not exceed 390 ohms (5 miles of 109-steel wire). The 390-ohm limit is determined by the talking battery supply requirements of the 500-type customer telephone sets when the battery is supplied from the remote P1 terminal power supply. The P1 carrier system has been designed to operate with the 500-type telephone set. The improved dialing, ringing and transmission features of that set will help to insure satisfactory performance of the

over-all carrier derived circuit. In keeping with the system objectives, the over-all transmission of a carrier channel and its voice frequency extension, using the 500-type telephone set, will be as good or better than that obtained on long rural lines using physical plant laid out by the Resistance Design Method mentioned earlier.

As shown in Fig. 27, the normal and staggered grouped frequency arrangements used for repeater operation allow 30-db bare line attenuation at the top frequency (96 kc) between the central office terminal and the first repeater or between repeaters, and about 30 db between the last repeater and each remote channel terminal at the top frequency used for that channel. Directional filter characteristics limit the repeater system can use a maximum of four repeaters for a total line loss of about 150 db at 96 kc. However, noise and crosstalk requirements will permit no more than two of the four repeaters to be used in the open-wire line, with the last cable repeater at least one mile back in the cable from the cable-open-wire junction, as shown in Fig. 27. Spacings must be limited to somewhat less than 30 db on certain line facilities such as B rural wire to insure proper terminal regulation.⁵



(a) STACKABLE FREQUENCY ARRANGEMENT: NON-REPEATED



NOTE 1
CHECK MAXIMUM LOSS AT 30KC AND IF LESS THAN GAIN IN REMOTE TERMINAL TO CENTRAL OFFICE DIRECTION, PLACE INPUT PAD EQUAL TO DIFFERENCE AT INPUT OF REPEATER.

NOTE 2
CHECK MINIMUM LOSS AT MINIMUM FREQUENCY OF EACH CHANNEL FOR THE LAST REPEATER TO REMOTE TERMINAL SECTION AND IF THIS IS LESS THAN THE REPEATER GAIN AT THAT FREQUENCY, PLACE PAD IN OUTPUT OF TERMINAL TO BUILD SECTION OUT TO REPEATER GAIN VALUE.

(b) GROUPED FREQUENCY ARRANGEMENTS: REPEATED

Fig. 27 — P1 carrier application ground rules.

⁵ C. C. Lawson, Rural Distribution Wire, Bell Lab. Record, pp. 167-170, May, 1954.

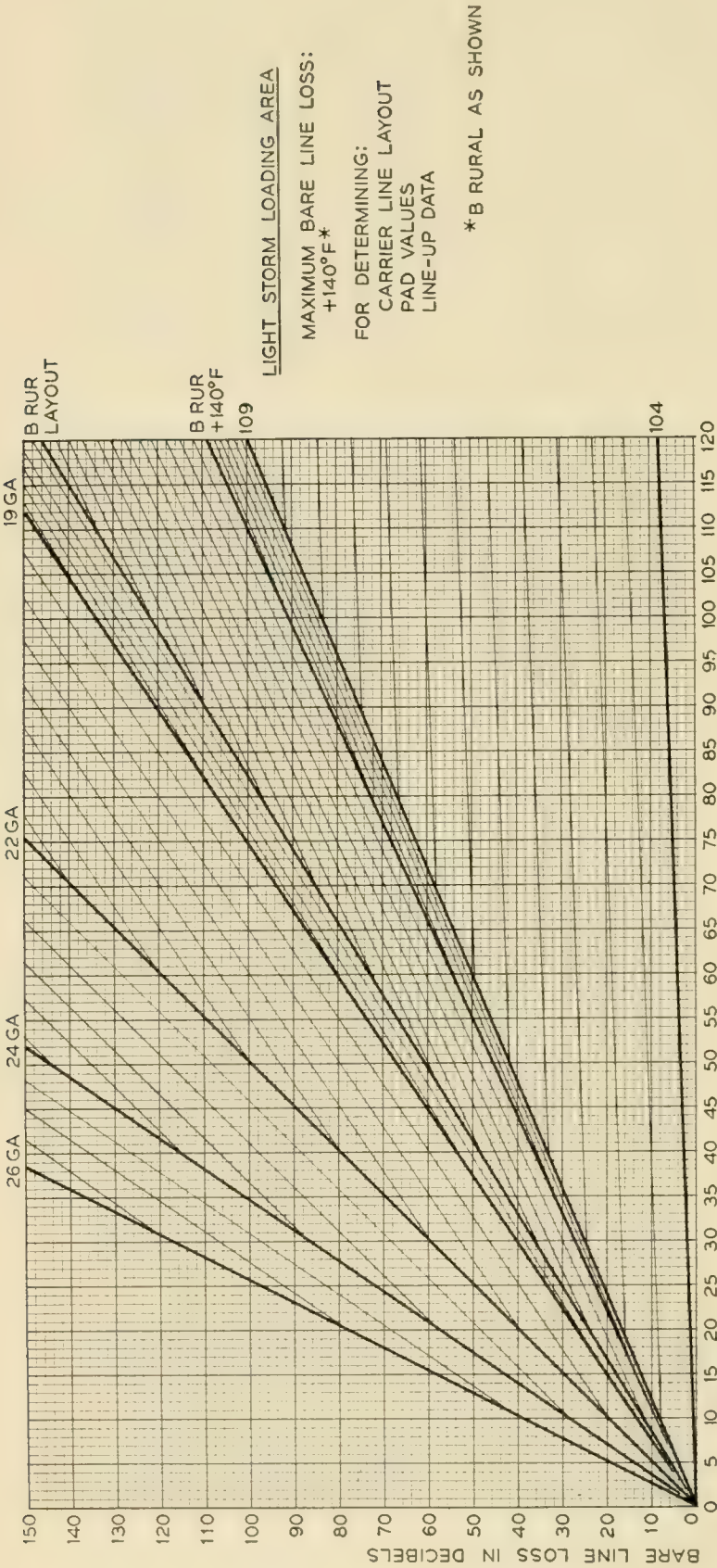
In addition to bare line loss, the ground rules make allowance for approximately 3 db of miscellaneous losses in any normal channel layout, including bridging losses of carrier blocking networks and other terminals on the carrier line, insertion losses of high-pass filters, and losses in the autotransformer and junction line filters used at the cable-open-wire junction. Since these losses do not all add directly, it is simpler to use an average loss factor to cover most conditions rather than make computations to determine a definite loss for each set of conditions that might exist. Thus for channels using a stackable frequency arrangement, a maximum of about 33 db loss, including the bare line loss and miscellaneous losses, may be expected between the points where the terminals connect to the line.

A further loss is experienced because the remote terminals are bridged onto the carrier line. As a result the carrier power transmitted toward the central office terminal is only +0.5 dbm due to a bridging loss of about 3.5 db at that point. Therefore, in the remote-to-central office direction, the minimum power will be -32.5 dbm (0.5 dbm - 33 dbm) at the line terminals of the central office terminal. The minimum carrier power in the central office to remote direction will be -29 dbm (+4 dbm - 33 dbm) at the bridging point of the remote terminal.

7.2 *Terminal and Repeater Location*

In laying out the carrier line design, it is first necessary to determine the possible locations for the remote terminals based on distances to the customers to be served and the availability of commercial ac power, since this is the most economical power source. (When commercial ac power can not readily be made available, the primary air cell batteries can be used.) Having determined the ideal location of the terminals from a physical standpoint, the makeup of the physical circuits back to the central office must be determined and computations made of the carrier frequency attenuation of the facilities. These loss computations are used to determine the number of repeaters required, if any, and their locations, once again modified by availability of commercial power. The Plant Engineer must also check for the necessity of input and output pads at the terminals and repeaters.

The need for loss computations led to the development of length-loss charts so that a carrier line design could be made in a manner very similar to the loop cable design using the Resistance Design Methods as mentioned earlier.⁴ Fig. 28 shows one of the 96-kc length-loss charts used to lay out repeater spacings and Channel 4 over-all circuit design. Fig. 29 shows the 48-kc length-loss charts as an example of the charts that are provided at each carrier frequency other than 96 kc for terminal-to-

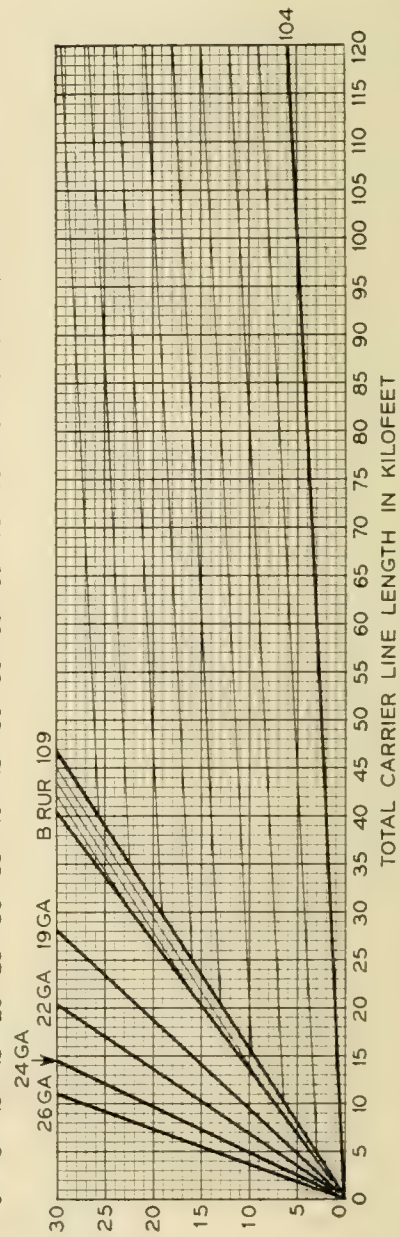


Light, Medium or Heavy Storm Loading Area

MINIMUM BARE LINE LOSS:
-40°F

FOR DETERMINING:
CARRIER LINE LAYOUT
PAD VALUES
LINE-UP DATA

LENGTH - DB LOSS CHARTS, 96 KC			
TYPE OF SYSTEM	CHAN NO.		CARRIER
	STACKABLE		UPPER
NORMAL	4	4	UPPER



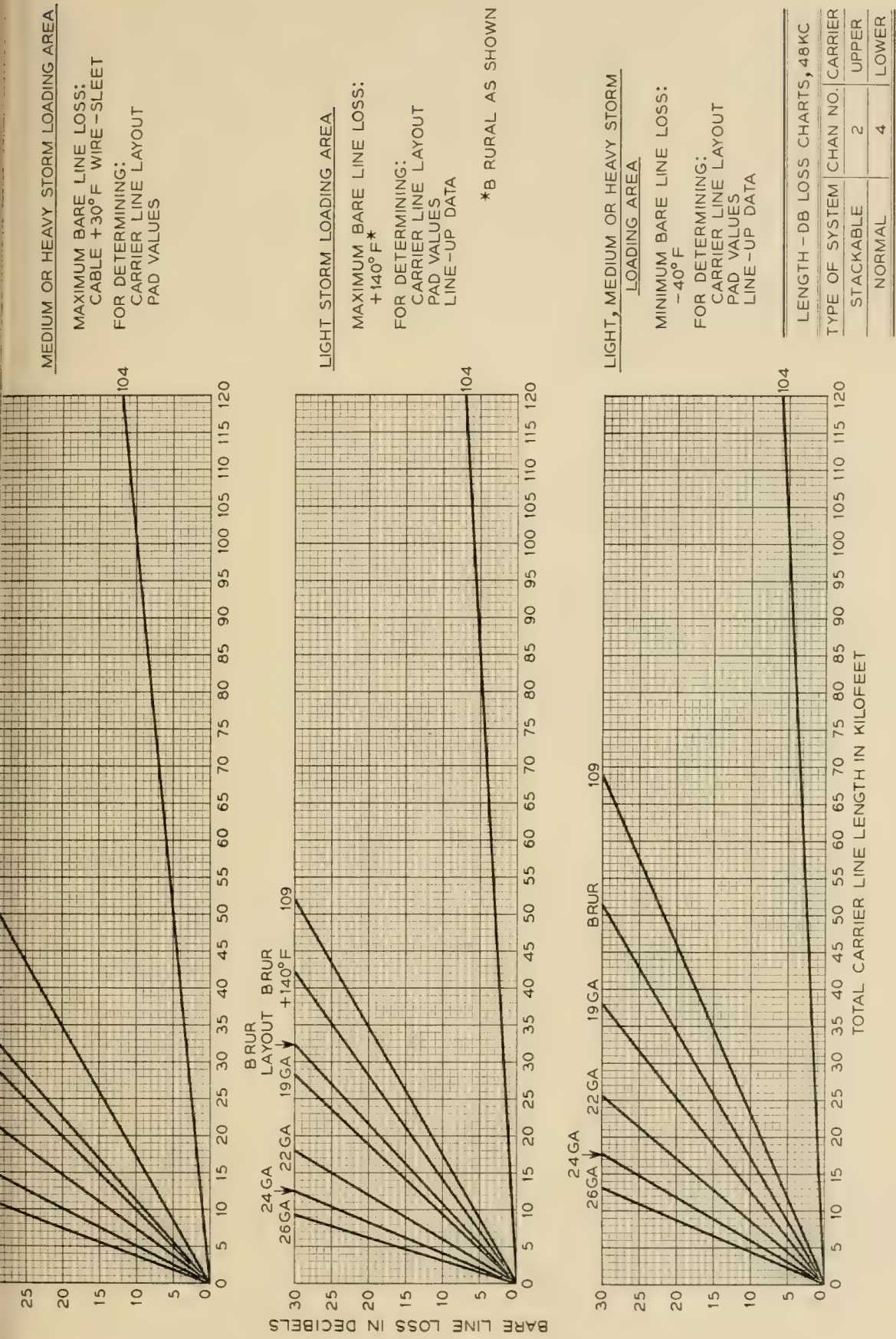


Fig. 29 — 48-ke length-loss charts.

the carrier assignments for the various frequency arrangements on a one and two-crossarm route using the R1 transposition design. Use of a transposition system giving better carrier frequency crosstalk performance than the R1 design is expected to permit the application of a number of additional channels over those shown in Fig. 2.

It will be noted that the grouped arrangement provides four channels less on a two-crossarm basis than the stackable arrangement. From a transmission standpoint, equal numbers of channels would be possible by assigning one or two channels to a pair, but these arrangements will usually be uneconomical for grouped systems because of the high cost per channel of repeaters.

7.5 *Line Networks*

The location of the line networks and filters, which are a permanent part of the carrier layout, will be designated by the Plant Engineer, and the location and type of the remaining networks, which will vary with changes in subscriber service, will be selected by the plant forces. The low-pass filters or carrier blocking networks used on the carrier lines are simple resonant circuits designed to match given ranges of capacitance that will be presented by the drop wire or open-wire side leads. Since this capacitance varies considerably with various lengths of facility, a method will be provided by which the total capacitance of the drop can be determined and the proper network chosen. The other line networks are applied to the line as necessary to achieve their particular functions.

8. INSTALLATION AND MAINTENANCE

A portable field test set has been developed which will simplify the installation and maintenance of the P1 carrier system. The new set, known as the 7F test set, will provide the carrier and audio frequencies and a means of measuring them required to align and troubleshoot units of the system. The set, which is battery operated, contains a carrier oscillator to supply test frequencies from 10 to 100 kc, an audio frequency oscillator having six selected frequencies in the range of 250 to 2,500 cycles, a modulator to modulate the carrier frequency signal with the audio signal, a demodulator for calibrating the modulated signals, and a wide-band amplifier-detector for making level and transmission measurements. The model of the set shown in Fig 31 included a precision dial for signaling testing which was subsequently found unnecessary and eliminated. An ac operated set providing the same desired facilities is now under development.

The carrier channel installation and lineup procedure is set up on the basis of using the test set and a generally available volt-ohmmeter to make a series of measurements in a specified order. This will permit potentiometers to be adjusted as necessary until the specified meter readings are obtained at built-in test points. Lineup of the terminals and repeaters done first in the central office to insure proper operation and then at the in-plant locations to check system performance.

Maintenance will be handled on a complete terminal replacement basis and will consist of making a series of checks with the test set to determine whether the terminal is functioning satisfactorily. If it is not and it cannot be adjusted to restore satisfactory operation, a replacement terminal will be used to restore service. All repairs and isolation of trouble within the terminal unit or on the individual boards will be handled at a centralized testing or repair point so as to require a minimum of personnel with electronic experience. The test set has been designed to handle all tests for a P1 carrier system, when used with the volt-ohmmeter, and when used by trained personnel will permit trouble to be isolated to a given printed wiring board in the P1 carrier equipment.

9. ACKNOWLEDGMENTS

The authors wish to thank all their colleagues for their important and necessary contributions to this paper. Particular appreciation should be expressed to E. H. Perkins for his contribution to the first half of the paper and to D. H. Smith for Section 6 on power supplies.

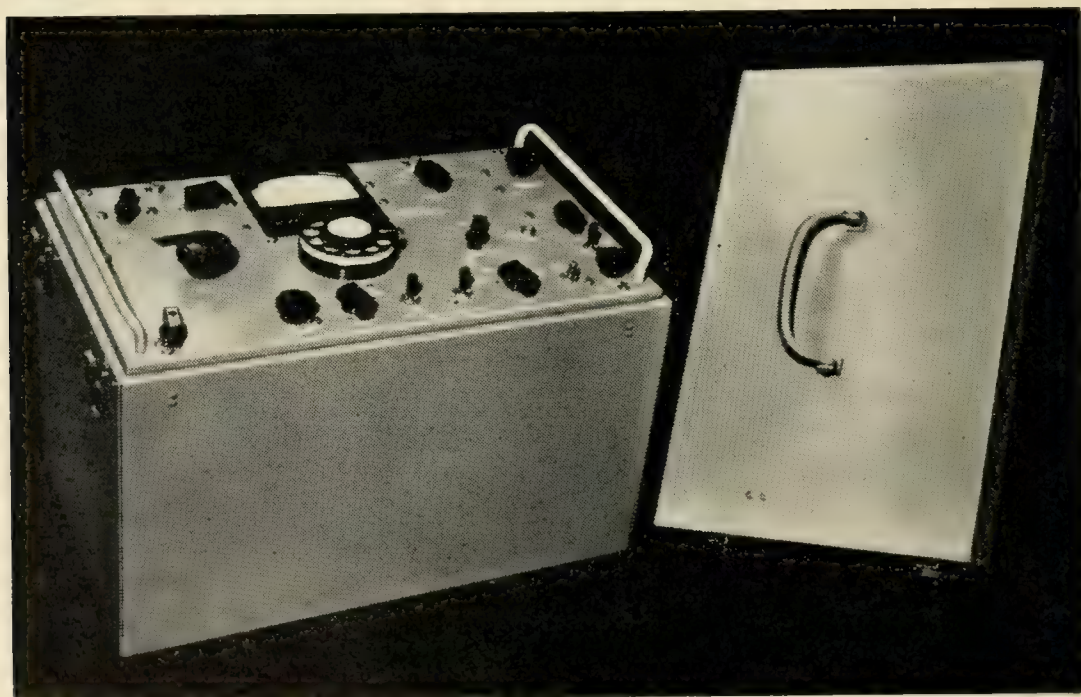


Fig. 31 — P1 carrier 7F test set.

An Experimental Dual Polarization Antenna Feed for Three Radio Relay Bands

By R. W. DAWSON

(Manuscript received April 26, 1956)

The fundamental problems associated with coupled-wave transducers which operate over a 3-to-1 frequency band have been explored and usable solutions found. The experimental models described are directed toward the broad objectives of feeding the horn-reflector antenna with two polarizations of waves in the 4-, 6- and 11-kmc radio relay bands.

INTRODUCTION

There are at least two communications problems which require frequency selective filters that operate in waveguides over an approximately 3-to-1 frequency interval: (1) channel-separation filters for a circular-electric waveguide system in which it is desirable to use the medium from perhaps 35 to 75 kmc,¹ and (2) band-separation networks needed for the horn reflector antenna that permits simultaneous transmission or reception in the 4, 6 and 11 kmc bands with both polarizations.^{2, 3} The research reported in this paper was directed at determining the capabilities of coupled-wave transducers for solving such problems. Experimental work was directed toward the second problem (above) because it is more immediate.

Fig. 1, which is a schematic representation of the feed array, comprised of three sets of directional couplers, shows that the 4-kmc bands are

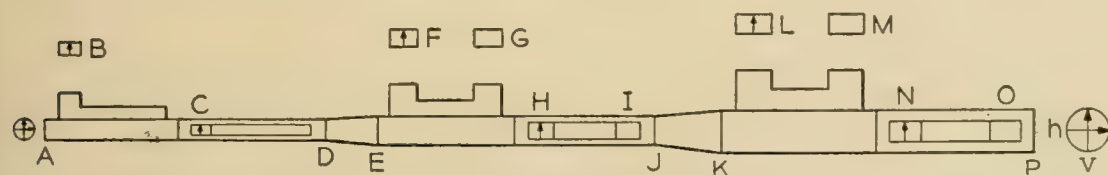


Fig. 1 — Schematic of dual-polarization feed for three microwave bands.

separated one at a time at the antenna end of the array. The 6-kmc bands are separated next and the 11-kmc bands are added or removed at the far end of the array.

GENERAL OBJECTIVES

Negligible loss of power should result when coupling $TE_{10}^{\square}$ waves to TE_{11}° waves for the six bands concerned.* The 6- and 11-kmc waves of both polarizations must pass through the round guide of the 4-kmc transducers without significant attenuation. Waves in the 11-kmc band must also pass through the circular guide of the 6-kmc transducers without appreciable loss. A good impedance match is desired at all ports. No cross coupling is desired between the orthogonally polarized waves in the round guide.

FUNDAMENTAL PROBLEMS ENCOUNTERED

The frequency-selectivity required to separate various bands in the same polarization can be achieved in a coupled-wave device by either varying the coupling coefficient and/or varying the phase constant, as illustrated by the expression for the amplitude of the selected wave:⁴

$$|E| = \frac{1}{\sqrt{1 + \left(\frac{\beta_1 - \beta_2}{2c}\right)^2}} \sin \left\{ \sqrt{1 + \left(\frac{\beta_1 - \beta_2}{2c}\right)^2} cx \right\} \quad (1)$$

where c = coupling coefficient

x = length of coupling array

β = phase constant

In the present designs some of the frequency selectivity is in the coupling holes. The greater part of the selectivity is in the design of the phase constants; they are made equal in the band to be selected ($\beta_1 - \beta_2 = 0$) and very unequal

$$\left(\frac{\beta_1 - \beta_2}{2c} \gg 2c\right)$$

in the frequency bands to be passed.

The size of the coupling hole must be controlled to avoid coupling hole resonance in any of the three bands that may be present. This problem is especially bothersome in the 4 kmc coupler where signals are present

* As used in this article, superscripts $\circ$ and $\square$ refer to round and rectangular waveguides, respectively.

in all three bands. To keep the coupler length within a reasonable size, the individual hole dimensions must be on the order of $\lambda_0/4$ (at 4 kmc), which *will* permit coupling hole resonance within the 3-to-1 frequency band. A further consideration is the selection of the hole shape to avoid perturbing the TE_{11}° wave that is orthogonal to the strongly coupled TE_{11}° wave.

Spacing of the coupling holes must *not* be $\lambda_g/2$ to avoid: (1) large reflections in the driven waveguide, and (2) large backward-travelling waves in the adjacent coupled waveguide. This requirement is easily met in the 11-kmc coupler where only one band is present; however, the presence of signals in two or three bands makes the non $\lambda_g/2$ spacing more difficult in the 6- and 4-kmc couplers.

Another phenomena of importance exists in coupled waveguides operating over an extended frequency range. A coupling aperture in the side wall (see Fig. 1) may interact with a high-order mode at the latter's cutoff frequency, resulting in a significant perturbation of the desired coupling. For example, at the frequency where TE_{21}° passes through cut off, the coupling between TE_{11}° and $\text{TE}_{10}^{\square}$ will be perturbed *if* the coupling hole is of sufficient size. Small coupling holes do not allow this perturbation to manifest itself. Coupling holes in a realistic design do become large enough to allow this effect to appear. Since dominant mode guides in the 4- and 6-kmc bands can support other modes in the higher frequency bands, considerable caution must be exercised in selecting the round guide sizes on this account alone. (The size of the round guide is determined also by the phase velocity in the rectangular guide).

SELECTION OF COUPLING APERTURES

A series of holes in either the narrow or broad side of the rectangular guide can, in principle, be used to achieve complete power transfer from $\text{TE}_{10}^{\square}$ waves to TE_{11}° waves. The specific consequences of coupling through holes located along the center line of the broad side will be considered first. (Off center holes are not of interest because they couple $\text{TE}_{10}^{\square}$ waves to both polarizations of TE_{11}° waves in a frequency-sensitive way.) The transverse magnetic field H_{ϕ} and the electric field E_{ρ} of the TE_{11}° waves can couple to $\text{TE}_{10}^{\square}$ waves. When two fields couple, the backward wave in the undriven guide can be greater than the forward wave in the same guide. To avoid this possibility, transverse slots can be used to prevent electric field coupling. The coupling of a transverse slot increases as the frequency is increased which suggests that 11 kmc signals be introduced at the position nearest to the antenna because the largest tolerable apertures for an 11-kmc coupler would not perturb 6- or 4-kmc

waves. Unfortunately, coupling to slots in this orientation is small, requiring several hundred for complete power transfer. Such a large number would make the coupler too long.

Coupling through holes in the center of the narrow wall of the rectangular waveguide as shown in Fig. 2 allows only the longitudinal magnetic field H_z to couple when the electric field of the TE_{11}° wave is parallel to the hole containing wall. No coupling exists between the $TE_{10}^{\square}$ waves and the TE_{11}° wave having an electric field perpendicular to the plane of the hole. The use of longitudinal slots where practicable minimizes perturbation of this wave. Since the desired $TE_{10}^{\square} - TE_{11}^{\circ}$ coupling decreases by 15 db from the 4 kmc to the 11 kmc bands (for $1.872 \times 0.872''$ and 2.2'' diam. guide), the layout of Fig. 1 suggests itself since some coupling discrimination is present for the higher frequency waves that pass through the lower frequency couplers.

DESIGN OF 11-KMC COUPLER

The objective of this design is to transfer all of the power from a dominant mode rectangular guide into one polarization of the TE_{11}° forward traveling wave in an adjacent circular guide. Fundamental coupled-wave theory⁴ shows that phase velocities must be matched in the two guides to achieve complete power transfer.

A standard rectangular guide size is selected and the round guide size that has the same phase constant is calculated for the center of the 1,000-mc wide band. An approximate total length is selected for the series of coupling holes that permits the holes to be spaced approximately $\lambda_g/4$ apart. The hole spacing is not critical although the non-directional properties of $\lambda_g/2$ spacing must be avoided. The required magnitude of multiple discrete couplings is shown in equation (40) of Reference 4 to be:

$$\alpha = \sin \left(\frac{\pi/2}{n} \right) \quad (2)$$

where n is the number of coupling holes and α is the amplitude of the wave transferred at a single coupling hole for unit incident amplitude.

Equation (3) expresses the power coupled from TE_{11}° waves to $TE_{10}^{\square}$ waves through a circular hole in a common wall of zero thickness where P_2 is the power propagating away from the coupling point in either direction in the undriven guide, and P_1 is in the driven guide. This derivation is based on the work of H. A. Bethe⁵ and some unpublished notes of S. P. Morgan.

$$\frac{P_2}{P_1} = \frac{0.6805\lambda_0^2 r^6}{ba^3 R^4 \sqrt{1 - \left(\frac{\lambda_0}{2a}\right)^2} \sqrt{1 - \left(\frac{\lambda_0}{3.413R}\right)^2}} \quad (3)$$

The quantities a and b are the large and small dimensions of the rectangular waveguide and R is the round guide radius. The wavelength in air is designated by λ_0 , and the radius of the coupling hole by r . A correction for the finite thickness of the wall is made by considering the circular coupling hole to be a round waveguide beyond cutoff.⁶ The additional loss is

$$\Delta = \frac{16t}{r} \sqrt{1 - \left(\frac{3.413r}{\lambda_0}\right)^2} \text{ (decibels)} \quad (4)$$

where t is the wall thickness. Total coupling loss per hole is defined by

$$20 \log_{10} \alpha = 10 \log \frac{P_2}{P_1} - \Delta \quad (5)$$

The number of coupling holes n is found from the approximate coupling length and hole spacing. Equation (2) is used to find α and then the hole radius r is calculated from (3).

Waveguide dimensions must be corrected to allow for the perturbation of the phase constants due to the coupling holes. The perturbed phase constant* for the round guide is

$$\beta_p^\circ = \beta^\circ + \frac{\sqrt{p^\circ}}{d} \quad (6)$$

where d is the hole spacing and p° is the coupling between a pair of round guides

$$p^\circ = \frac{0.1056r^6\lambda_0^2}{R^8 \left[1 - \left(\frac{\lambda_0}{3.413R}\right)^2 \right]} \quad (7)$$

The perturbed phase constant for the rectangular guide is

$$\beta_p^\square = \beta^\square + \frac{\sqrt{p^\square}}{d} \quad (8)$$

$$p^\square = \frac{4\pi^2 r^6 \lambda_0^2}{9a^6 b^2 \left[1 - \left(\frac{\lambda_0}{2a}\right)^2 \right]} \quad (9)$$

* This correction is due to S. A. Schelkunoff as noted on page 708 in Reference 4.

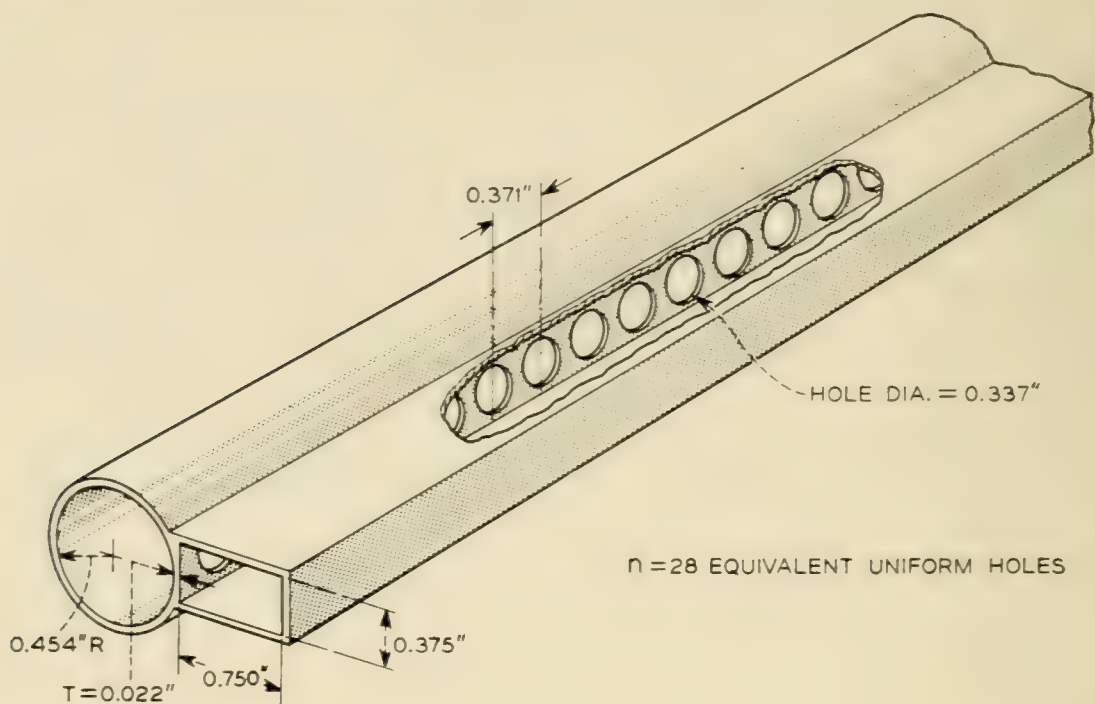


Fig. 2 — 11-kmc coupler sketch.

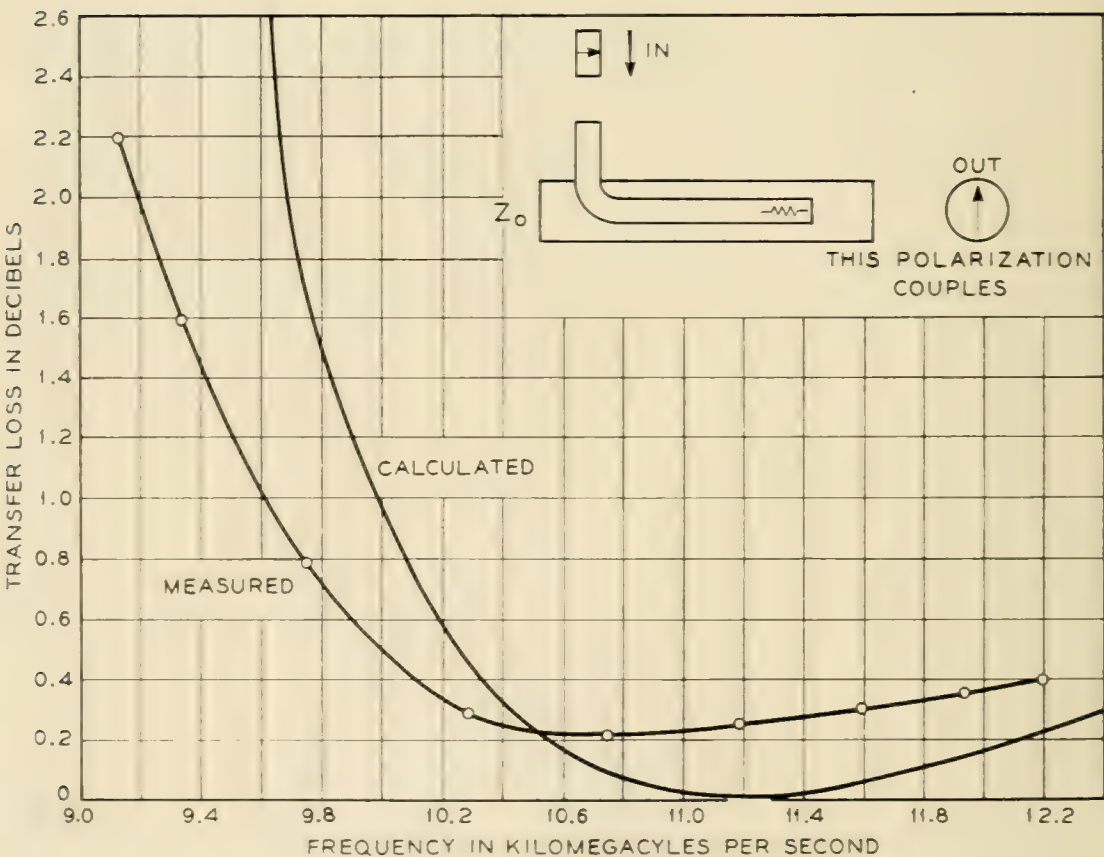


Fig. 3 — 11-kmc coupler transfer loss.

The perturbed phase constants are made equal through a suitable choice of R and a , this choice being somewhat influenced by the coupling-hole radius r .

Three coupling apertures of successively reduced size are used at the ends of the array of identical holes to produce four reflections having the relative amplitudes of 1, 2.7, 2.7, 1. The modified binomial distribution was chosen because impedance matching can be secured over a broader

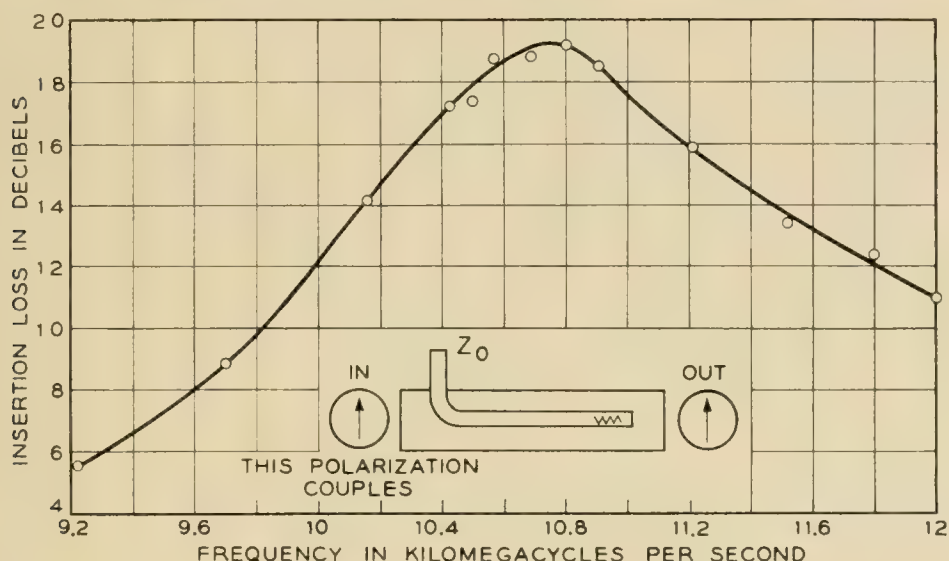


Fig. 4 — 11-kmc coupler insertion loss.

band (with minor degradation of the center-frequency match) when compared to a standard binomial distribution. Amplitude reflections* from the start of the coupling array are

$$A_r = \frac{\lambda_g Q}{4\pi d} \quad (10)$$

where Q is the reflection from a single coupling hole.

Fig. 2 is a sketch of the coupler with the final design dimensions. Fig. 3 shows the measured and theoretical transfer loss of an 11 kmc coupler. Fig. 4 indicates the measured insertion loss for the same coupler.

DESIGN OF 6-KMC COUPLER

The 6-kmc coupler as shown in Fig. 5 utilizes a partially dielectric-filled rectangular guide coupled to the circular guide. The use of dielectric loading makes it possible for the phase velocities to be equal in the two guides in the center of the 500-mc wide 6-kmc band, and unequal in

* Information given to S. E. Miller by S. A. Schelkunoff in an informal communication.

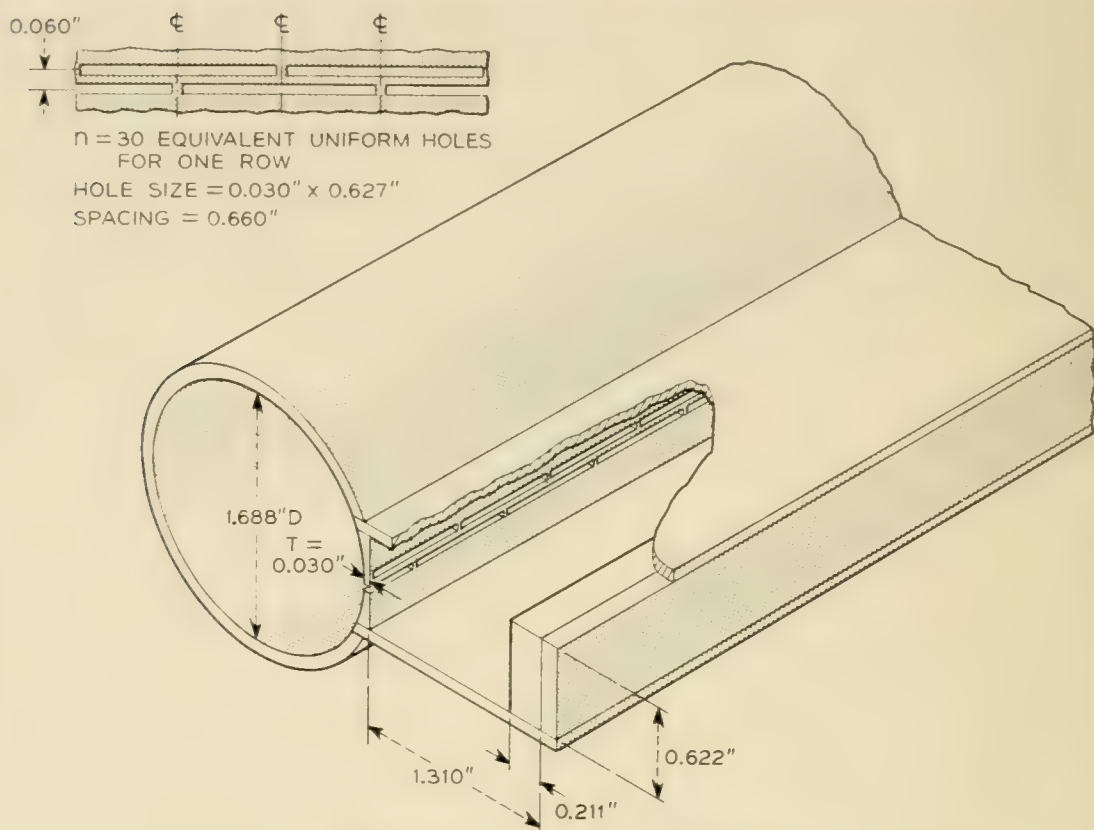


Fig. 5 — 6-kmc coupler sketch.

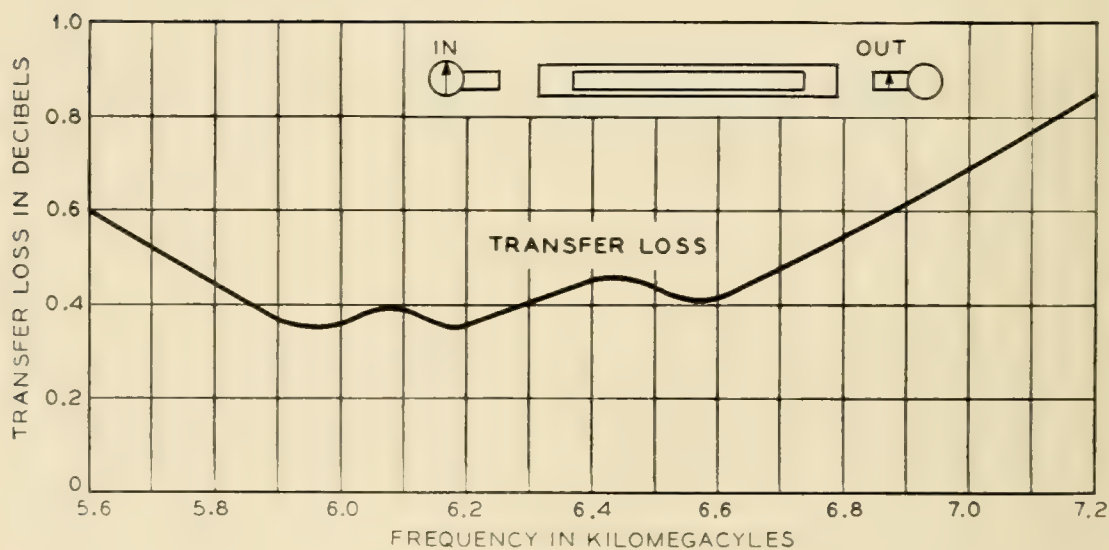


Fig. 6 — Transfer loss of 6-kmc coupler.

the 11-kmc band; thereby low transfer loss is obtained in the 6-kmc band and a high transfer loss in the 11-kmc band. Measurements have shown that when the cut-off frequencies of higher modes occur in the band of interest an uncontrolled increase of coupling may result. Special precautions are required in selecting the round guide size to avoid this condition. The design process is shown in Appendix I. Figs. 6 and 7 show the transfer and insertion losses in the 6-kmc band.

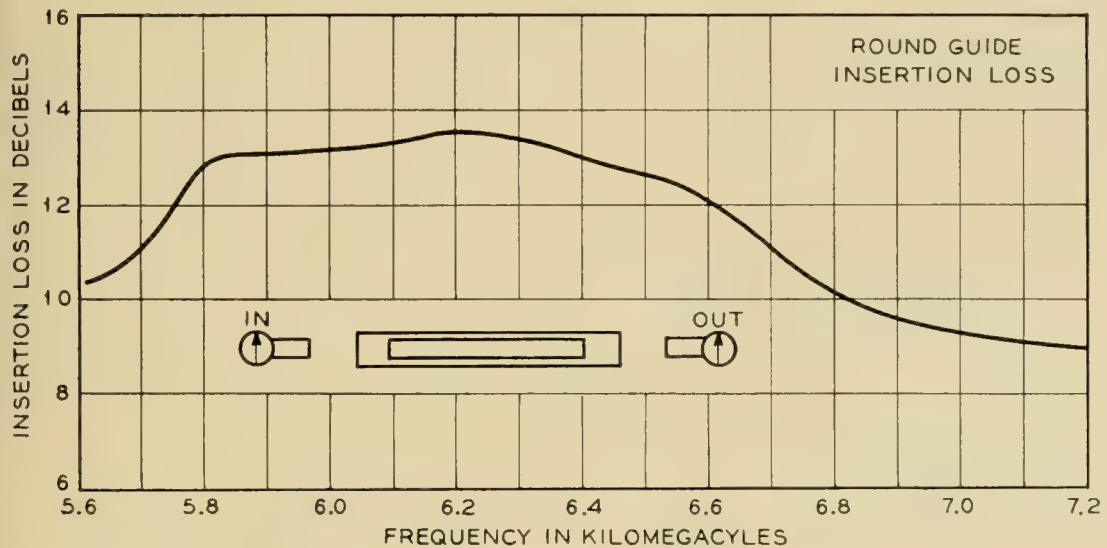


Fig. 7 — Insertion loss of 6-kmc coupler.

SLOT SIZE = 0.740" x 0.0786"
 SPACING = 0.780"
 POST DIA = 0.300"
 GUIDE SIZE = 1.724" x 0.872"

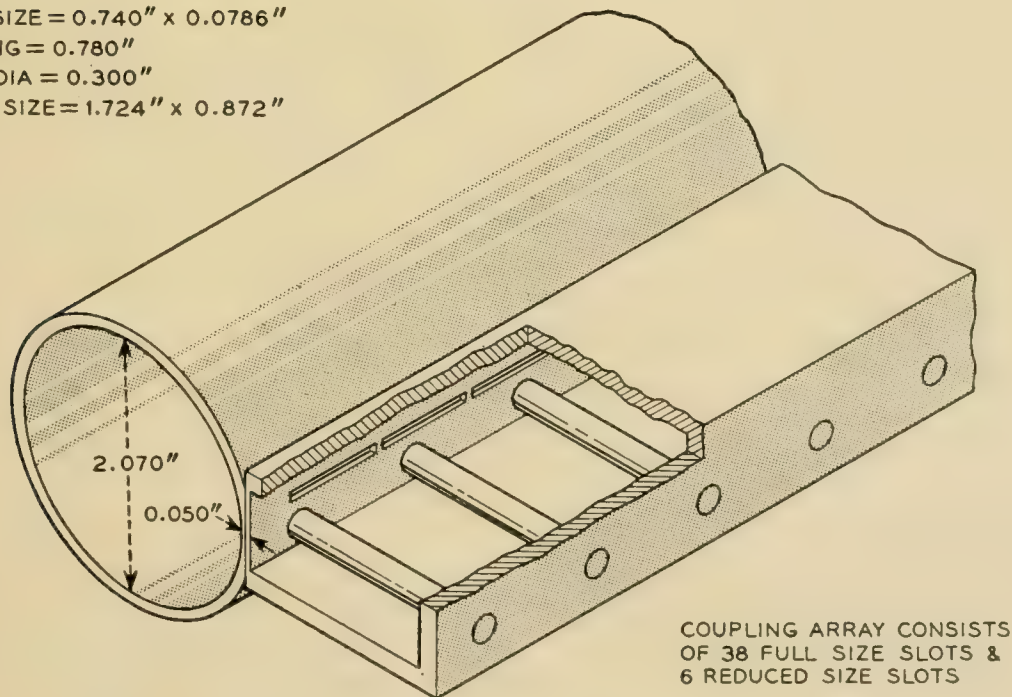


Fig. 8 — 4-kmc coupler sketch.

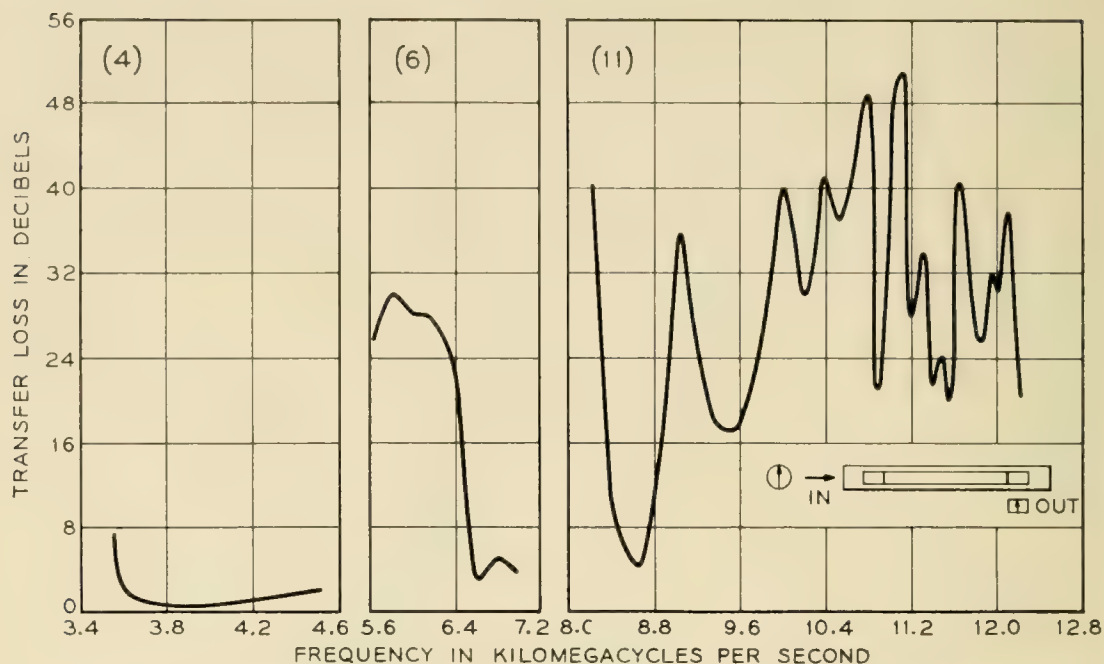


Fig. 9 — Transfer loss of 4-kmc coupler in 4-, 6- and 11-kmc bands.

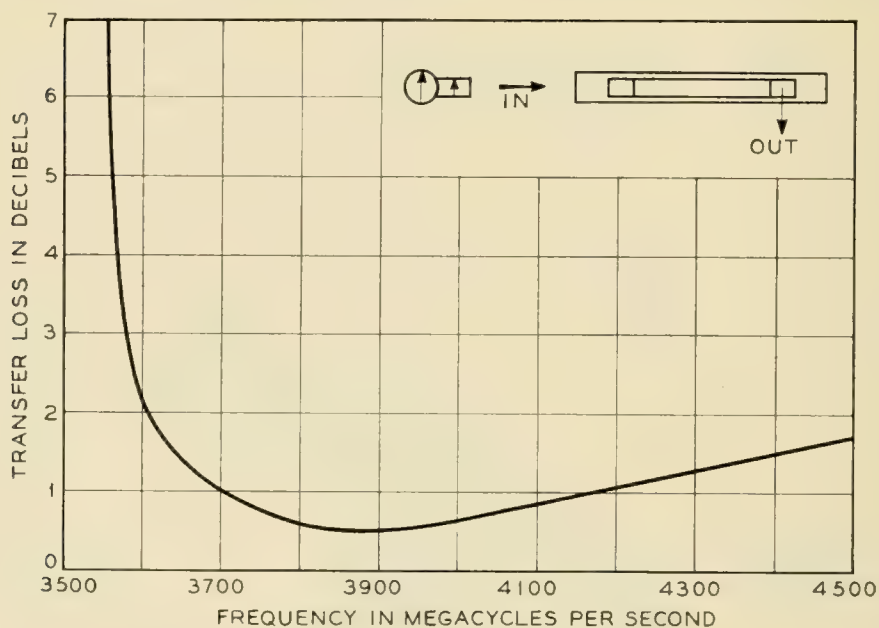


Fig. 10 — Transfer loss of 4-kmc coupler in 4-kmc band.

4-KMC DESIGN

The use of periodic loading in the rectangular guide is not suitable for use in a 4-kmc coupler design. When the phase constants are made equal in the 4-kmc band, the resulting difference of phase constants in the 6-kmc band is too small to create a sufficiently high transfer loss in that band. Periodic loading can produce the desired result.

Capacitive rods form a periodic structure in the rectangular guide, as shown in Fig. 8, that creates a rejection band in the 6-kmc region. Fig. 9 illustrates how effectively the rejection band increases the transfer loss in the 6-kmc region. Phase velocities are made equal in the rectangular and round guides at the center of the 500-mc wide 4-kmc band to secure a low transfer loss. Due to the rejection bands and the difference of phase constants, high transfer losses result in the 6- and 11-kmc bands. To prevent uncontrolled coupling the round guide size is chosen so that no modes cut off in the three bands. Details of the design are covered in Appendix II. Figs. 10 and 11 show the measured transfer and insertion losses in the 4-kmc band.

MEASURED CHARACTERISTICS OF ARRAY

The three pairs of couplers were assembled in a tandem array with linear taper sections between them. In the discussions that follow the port designations of Fig. 1 will be used. Transfer loss measurements indicate

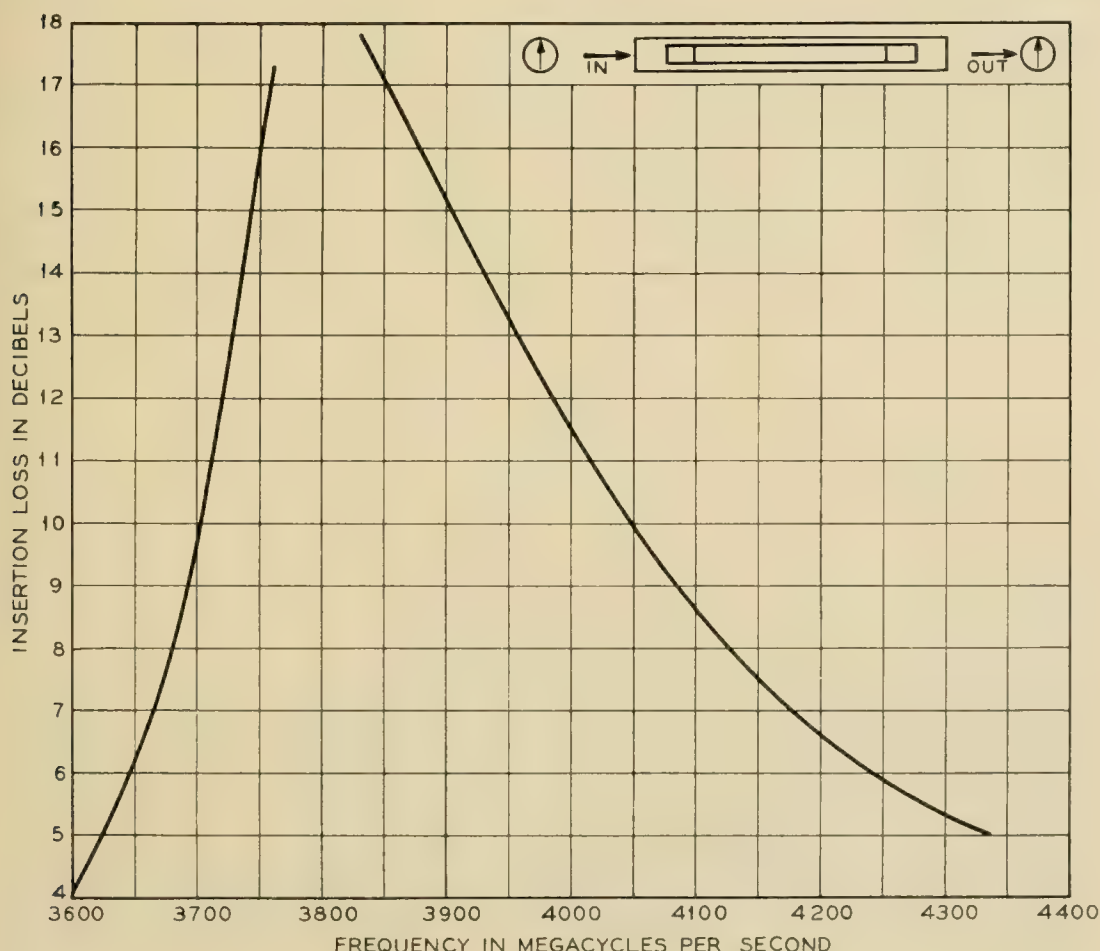


Fig. 11 — Insertion loss of 4-kmc coupler in 4-kmc band.

how much power in a forward traveling TE_{11}° wave is transferred to a forward traveling $TE_{10}^{\square}$ wave. Figure 12 shows that the coupling polarization transfer loss remains under 1.1 db in the three bands except for a small region in the 11-kmc band, while the transfer loss for the non-coupling polarization exceeds 20 db in the three regions as shown in Fig. 13. The return loss at Port P exceeded 23 db over the 4-kmc band. This result included the total reflection of 4-kmc signals from Taper $J-K$ after attenuation by twice the coupler insertion loss, and also included the reflections from the rectangular guide port N or L which

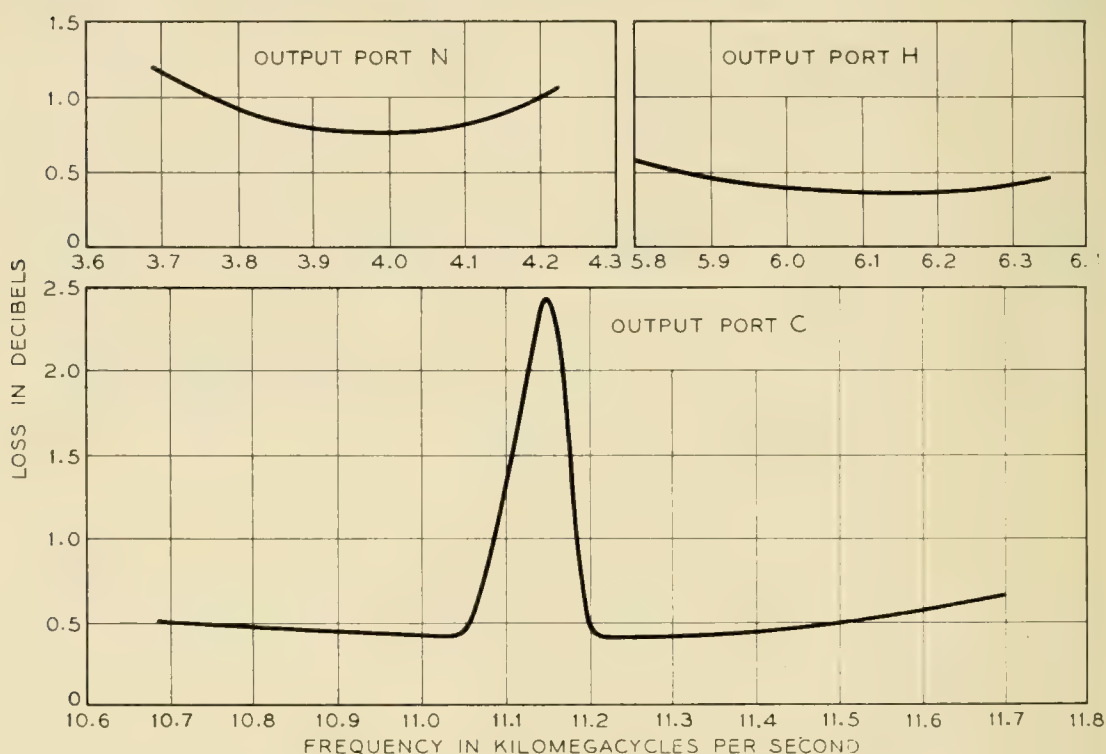


Fig. 12 — Transfer losses in the array for coupling polarization.

are separated from P by only the small transfer loss. Return loss for the 6- and 11-kmc bands exceeded 23 db at Port P . Cross polarization is the ratio of the energy in the coupling polarization waves to the orthogonal non-coupling polarization waves emerging at Port P . Cross polarization figures are no lower than 20, 32 and 22 db in the 4-, 6- and 11-kmc bands.

CONSTRUCTION OF COUPLERS

The coupler design requires that the coupling aperture exist in a narrow wall of the rectangular guide that is common to the round guide. The 4-kmc coupler consists of machined rectangular and round sections. A two-piece rectangular guide was milled from brass and

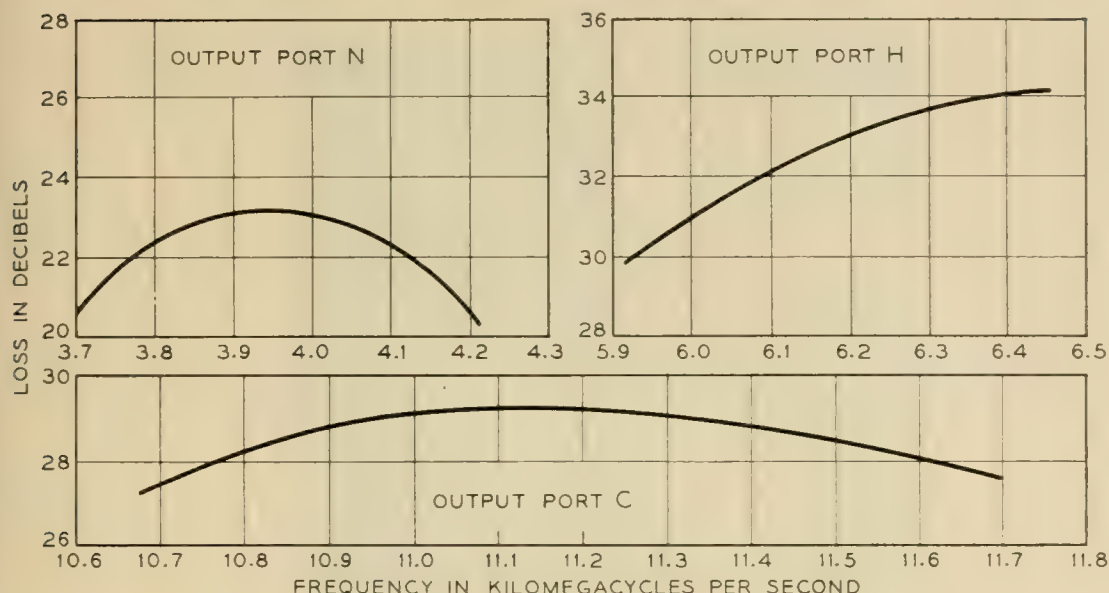


Fig. 13 — Transfer losses in the array for non-coupling polarization.

soldered to the round guide. The coupling slots were then cut through the common wall. An electroforming technique was used to fabricate the 6- and 11-kmc couplers. A rectangular guide with precut coupling holes was clamped to a round mandrel and the entire structure was electroformed. The mandrel was later removed by dissolving it in a hot concentrated solution of sodium hydroxide. A typical mandrel and rectangular guide is shown in Fig. 14. An illustration of the entire ensemble appears as Fig. 15.

DESIGN REFINEMENTS

Return loss at the round Port P might be improved in a revised design by broad-banding the $TE_{10}^{\square}$ - TE_{11}° transfer loss, with an associated increase in the TE_{11}° insertion loss. This might be done without introducing mode troubles, by using ridged waveguide. In Fig. 12 the abrupt peak of transfer loss for coupling polarization waves in the 11

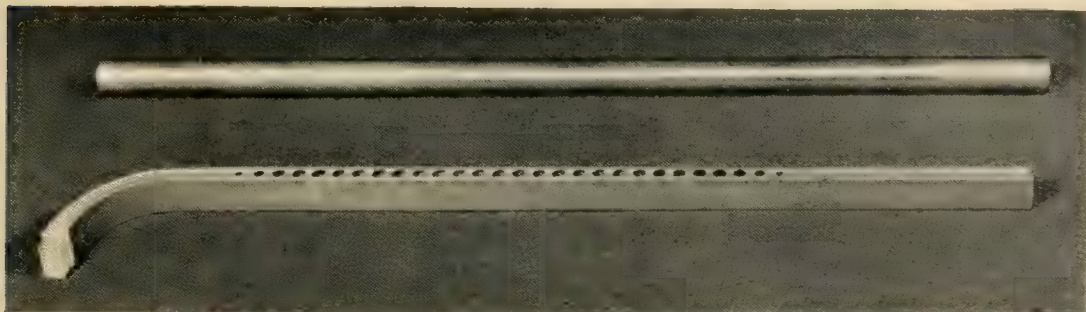


Fig. 14 — Mandrel and rectangular guide.

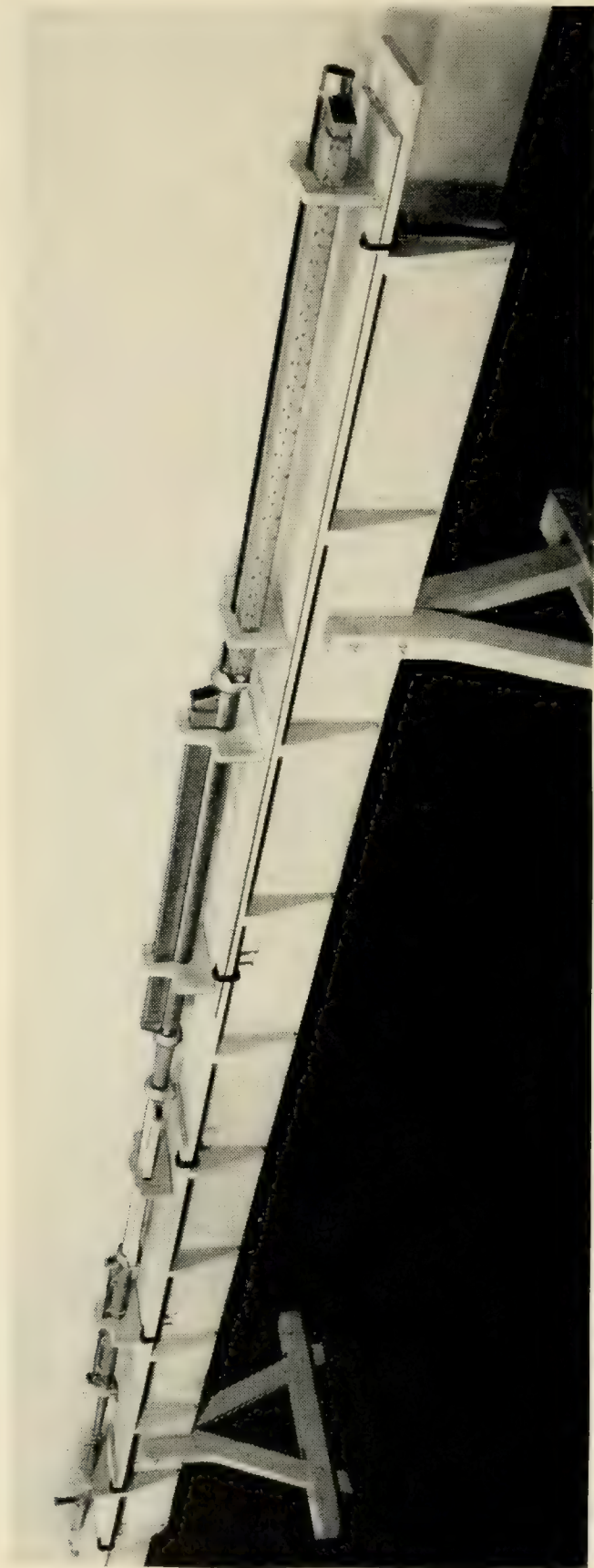


Fig. 15 — The array.

kmc band is due to an appreciable coupling of TE_{11}° waves to $TE_{30}^{\square}$ waves in the 4 kmc couplers in the vicinity of the $TE_{30}^{\square}$ cutoff. The $TE_{30}^{\square}$ cutoff might be moved above the 11-kmc band by using ridged rectangular guide.

SUMMARY

A combination of three pairs of coupled-wave transducers has successfully permitted six distinct bands to be fed from individual $TE_{10}^{\square}$ rectangular waveguides into the two orthogonal polarizations of TE_{11}° waves in a multi-mode round waveguide. The resulting structure enabled low transfer loss values to be obtained simultaneously over two 500-mc wide bands and one band 1000-mc wide distributed over a 3-to-1 frequency interval.

ACKNOWLEDGMENT

A substantial portion of this work was carried out as a joint project with H. E. Heskett and S. E. Miller.

APPENDIX I

DESIGN OF 6-KMC COUPLER

A rectangular guide size is chosen and a reasonable value selected for the phase constant such as $\beta^{\square} = 1.5\pi/\lambda_0$. The resulting round guide size must prevent modes from cutting-off within the 6-kmc band and preferably the cut-off frequency should be above the 6-kmc band. The thickness of the dielectric (polystyrene) strip is determined⁷ from (11), (12), and (13) where K_1 and K_2 are transverse wave numbers of the air-filled and of the dielectric sections of the guide.

$$K_1 = \sqrt{\left(\frac{2\pi}{\lambda_0}\right)^2 - \beta_{\square}^2} \quad (11)$$

$$K_2 = \sqrt{\epsilon_r \left(\frac{2\pi}{\lambda_0}\right)^2 - \beta_{\square}^2} \quad (12)$$

$$\cot K_2 d = -\frac{K_1}{K_2} \cot K_1(a - d) \quad (13)$$

In the above equations ϵ_r is the relative dielectric constant and d is the slab thickness. After solving for d the equations are resolved for K_1 and K_2 values at the 6-kmc band edges and also at the lower end of the 11-kmc band. Resulting rectangular-guide phase constants should cause a very

low and a high value of transfer loss as indicated by (14) which represents the forward traveling coupled wave amplitude where cx is the total coupling strength.⁴

$$E_2 = \frac{1}{\sqrt{\frac{(\beta_{\square} - \beta_{\circ})^2}{4c^2} + 1}} i \sin \left[\sqrt{\frac{(\beta_{\square} - \beta_{\circ})^2}{4c^2} + 1} \right] cx \quad (14)$$

At midband $cx = \pi/2$ and experience has shown $x \cong 10 \lambda_0$; therefore $c \cong \pi/20\lambda_0$ at midband. The coupling hole radius is found from (3). Because longitudinal slots are used instead of round holes the equivalent hole radius r is found from $\frac{4}{3}r^3 = P\ell^3$, where P is the magnetic polarizability⁵ and ℓ is the length of the chosen slot. To avoid slot resonance in either band, the length was chosen to be approximately $\lambda_0/4$ in the 6-kmc band and $\frac{5}{8}\lambda_0$ in the 11-kmc band. The power expression of (3) must be corrected by the wall thickness effect (4) and also multiplied by the factor F due to the presence of the dielectric slab⁷ (15).

$$F = \frac{2a^3 K_1^4}{\pi \beta_{\square} \lambda_{\square} [K_1(\theta_1 - \frac{1}{2} \sin 2\theta_1) + A^2 K_2(\theta_2 - \frac{1}{2} \sin 2\theta_2)]} \quad (15)$$

$$\theta_1 = K_1(a - d) \quad \theta_2 = k_2 d \quad A = \left(\frac{K_1}{K_2} \right)^2 \frac{\cos \theta_1}{\cos \theta_2}$$

Theoretical coupling loss per hole is defined by

$$20 \log_{10} \alpha = 10 \log \frac{P_2}{P_1} - (\Delta + 10 \log_{10} F) \quad (16)$$

An additional correction which reduces the coupling loss is due to the long length of the slot. Although the slot resonates near 9 kmc, an increase of 3 db in a single slot coupling results at 6 kmc. This effect was found experimentally from a sample test line with several slots. To avoid excessive length two rows of coupling slots were employed. They were staggered to improve the continuity of coupling from discrete points. An approximate design is on hand at this point. The final dimensions for the guides and coupling holes are found after the perturbations of the phase constants are considered by the same process as noted in the discussion of the 11-kmc design. Impedance matching for the dielectric strip and the coupling slot array was patterned after the technique shown for the coupling hole array in the 11 kmc design.

APPENDIX II

DESIGN OF 4-KMC COUPLER

A round guide size is selected so that no modes are cut-off in the three bands. Coupling power varies with wavelength as shown in (17) for TE_{10} to TE_{11}° coupling.

$$Y = \frac{\lambda_0^2}{\left[1 - \left(\frac{\lambda_0}{3.413R}\right)^2\right]} \quad (17)$$

To maintain the minimum variation across the band, the following requirements must be met which were deduced from the coupled wave theory⁴ and (17).

$$\sin (cx)_1 = \sin (cx)_2 \quad (18)$$

$$\frac{(cx)_1}{(cx)_2} = \frac{\lambda_0}{\sqrt{1 - \left(\frac{\lambda_0}{3.413R}\right)^2}} \quad (19)$$

$$(cx)_1 + (cx)_2 = \pi \quad (20)$$

The equations are solved for $\sin cx$ (the minimum band edge transfer loss) which for a diameter of 2.10" is 0.5 db. A value of 30 db is chosen for α , (2), as the first approximation. Because the band edge transfer loss is 0.5 db due to frequency variation of coupling, an additional loss of only 0.2 db is allowed for the phase constant difference $(\beta_L^{\square} - \beta^{\circ})$. It is now necessary to make $\beta^{\circ} = \beta_L^{\square}$, where $\beta_L^{\square}$ is the phase constant of the loaded rectangular guide. Equation (21) gives a theoretical periodic loading formula where L is the spacing between loading elements.⁹

$$\begin{aligned} \cos \beta_L^{\square} L &= A \cos (\beta^{\square} L + \Phi) \\ \Phi &= \arctan \frac{b_0}{2} A = \sqrt{1 + \left(\frac{b_0}{2}\right)^2} \end{aligned} \quad (21)$$

Assume initially that $L = \pi/\beta^{\circ}$. The required susceptance b_0 of the capacitive rods can be found experimentally from a loaded test line by varying b_0 until the first rejection band covers the 6 kmc band. An iterated process is used to find L because it is dependent on $\beta^{\square}$ which is the parameter being sought. Measurements indicated that (21) does not predict $\beta^{\square}$ very accurately and for that reason an experimental adjustment of $(\beta_L^{\square} - \beta^{\circ})$ is desirable.

The guide dimensions are now known; however, they must be corrected for the perturbations of the phase velocities as outlined in the 11-

kmc coupler design. A single row of longitudinal coupling slots which are not resonant in the 6- or 11-kmc bands is used. The radius of a round hole equivalent to the slot is obtained as in Appendix I, by setting $\frac{4}{3}r^3 = Pl^3$.

Impedance matching of the coupling array and of the loading elements is accomplished by tapering the amplitude of the end elements as indicated by a discrimination function of the coupled wave theory.⁴ For a 5 element series $1 - n_1 - n_2 - n_1 - 1$

$$D = \frac{2(n_1 + 1) + n_2}{2 \left(n_1 \cos \frac{4\pi Z}{\lambda_g} + \cos \frac{8\pi Z}{\lambda_g} \right) + n_2} \quad (22)$$

where Z is the spacing between elements.

The discrimination D is set equal to infinity which permits the denominator to be set equal to zero and solved simultaneously for n_1 and n_2 by using both band edge wave-lengths. Round hole sizes are readily obtained since the coupling coefficient is directly proportional to the cube of the hole radius from which the necessary equivalent longitudinal slot can be calculated. Susceptance values of the capacitive posts are found from the absolute value of the reflection coefficient which equals

$$\frac{b_0}{\sqrt{b_0^2 + 4}}$$

REFERENCES

1. S. E. Miller, Waveguide as a Communication Medium, B.S.T.J., Nov., 1954.
2. A. T. Corbin and A. S. May, Broadband Horn Reflector Antenna, Bell Laboratories Record, **33**, p. 401, Nov., 1955.
3. A. P. King, Dominant Wave Transmission Characteristics of a Multimode Round Waveguide, Proc. I.R.E., **40**, Aug., 1952.
4. S. E. Miller, Coupled Wave Theory and Waveguide Applications, B.S.T.J., May, 1954. (See page 681.)
5. H. A. Bethe, Physical Review, **66**, p. 63, 1944. Also Report 43-22, Lumped Constants for Small Irises, Mar. 24, 1943; Report 43-26, Formal Theory of Wave Guides of Arbitrary Cross Section, Mar. 16, 1943; and Report 43-27, Theory of Side Windows in Wave Guides, Apr. 4, 1943; from M.I.T. Radiation Lab.
6. N. Marcuvitz, Waveguide Handbook, p. 408, McGraw Hill.
7. H. Seidel, private communication re: Slab Width Determination and Effects on Coupling Parameters in Partial Dielectric Loading.
8. S. B. Cohn, Determination of Aperture Parameters by Electrolytic Tank Measurements, Proc. I.R.E., Nov., 1951.
9. J. C. Slater, Microwave Electronics, p. 183, Van Nostrand.

The Character of Waveguide Modes in Gyromagnetic Media

By H. SEIDEL

(Manuscript received August 31, 1956)

A magnetized gyromagnetic medium is birefringent. The effect of birefringence is studied in rectangular and circular waveguides with special attention paid to propagation characteristics in guides of arbitrarily small cross-section. Propagating, small-size structures are found in certain ranges of magnetization for both types of guide.

I. INTRODUCTION

A gyromagnetic medium, isotropic in the absence of a magnetizing field, becomes axially symmetric with respect to that field when magnetized. A tensor susceptibility¹ is thus produced which reflects the resulting anisotropy. Two essentially different types of rays appear in the medium in much the same manner in which the ordinary and extraordinary optical rays form in a calcite crystal. These rays may combine to produce results in a ferrite loaded waveguide quite alien in character to those of a conventional isotropic guide. Since the ferrite is, to first order, characteristic of general gyromagnetic media we shall discuss all gyromagnetic phenomena in terms of ferrites alone.

One very startling phenomenon observed in ferrite loaded waveguides is the occurrence of propagation in a waveguide of arbitrarily small transverse dimensions.² We shall show that this type of wave guide behavior is a consequence of the particular form of the birefringent character of the medium.

In order to understand the nature of the ferrite loaded case let us first consider the conventional isotropic small wave guide. Fig. 1 shows, schematically, the field distribution encountered in a small rectangular waveguide operating in a (1,1) mode. The x axis is shown along the wide transverse dimension and z is along the narrow height dimension. The y axis is chosen to coincide with the guide axis.

The field solutions of such a waveguide may be obtained as a super-

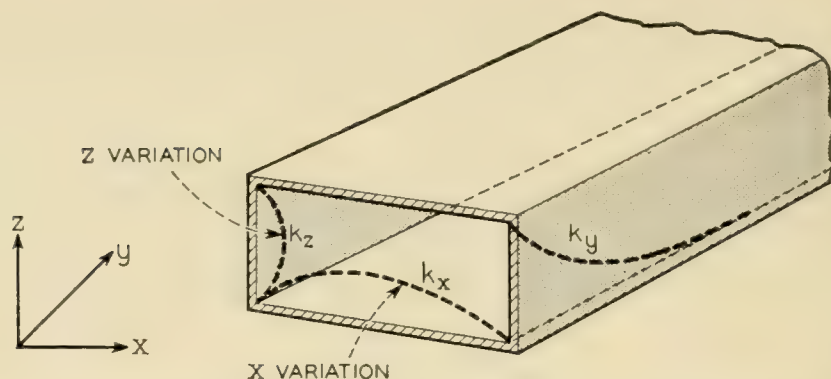


Fig. 1 — Rectangular waveguide mode in isotropic medium for cutoff guide.

position of plane waves of dependence $\varepsilon^{-i\mathbf{k}\cdot\mathbf{R}}$. If we represent $\mathbf{k}$ in a cartesian frame, the wave equation is satisfied for the condition

$$k^2 = k_x^2 + k_y^2 + k_z^2 = \omega^2 \mu \varepsilon$$

where μ and ε are the permeability and permittivity respectively of the medium. Satisfaction of wall boundary condition requires that k_x and k_z be real and that each be of the order of the reciprocal of the transverse guide dimensions. Small transverse dimensions thus cause k_y^2 to be negative, driving the waveguide into a cutoff condition.

We shall now find that birefringence permits another class of modes in the small size ferrite loaded waveguide. Letting the magnetic axis be in the z direction, it will be shown in the text that corresponding to any mode of the guide k_y and k_z are unique. Birefringence generally requires that two different magnitudes of $\mathbf{k}$ occur simultaneously, causing two different values of k_x to appear. In particular, let us postulate that both these values of k_x are imaginary. Given two exponentials, it is possible now to satisfy the requirements of electric field nulls at either side wall, as shown in Fig. 2. At the other side wall we shall show that the exponentials decay so fast as to effectively cause the field to vanish there. Since $k_{x1,2}^2$ are now negative quantities, there is no contradiction in presuming that k_y^2 may now be positive, thus permitting propagation in an arbitrarily small size waveguide.

The effect of birefringence may then be that of transforming a class of longitudinally cutoff modes into another class that propagates longitudinally but cuts off transversely. The condition of this occurrence will be shown to be that for which the diagonal term of the Polder tensor, μ , is positive and is less in magnitude than the magnitude of the off diagonal term κ . In the case of a small rectangular guide, propagation occurs anomalously for negative values of μ , as well but in a manner not as

substantially dependent on the birefringent character of the medium for large width to height aspect ratios of the waveguide. We shall find, further, that propagation occurs with entirely real values of k_x and k_z .

It will be shown that the proper wave equation for one of the two birefringent rays is satisfied in the small waveguide limit by the relationship

$$k_x^2 + k_y^2 + k_z^2/\mu = 0.$$

In the region of $\mu > 0$, and k_z real, we confirm somewhat more rigorously the requirement stated earlier that either k_x or k_y be imaginary. However, k_x and k_y may both be real over a range of negative values of μ , permitting boundary conditions to be satisfied, approximately, in waveguides having aspect ratios of the type discussed earlier, by just one class of rays in the small size waveguide.

Propagation in small size circular guide employing the essential character of birefringence, occurs over the entire range of $|\mu| < |\kappa|$. This range is divided into that of $\mu > 0$ and that of $\mu < 0$. Transmission occurs in one sense of circular polarization in each of these regions and for both senses for $\mu < 0$. Thompson³ has suggested that propagation in a small circular waveguide might be attributed to the negative permeability of one preferred polarization; it appears, however, that propagation is possible over a considerably wider range of conditions and for somewhat different reasons.

In the case shown in Fig. 2, higher propagating modes occur in a rectangular waveguide when one half or more sinusoids of field variation occurs in the z direction. These simply produce the result of stronger

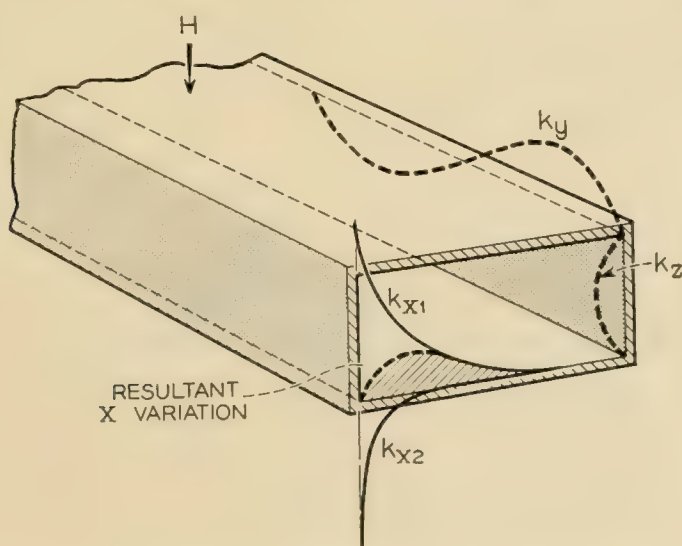


Fig. 2 — Mode in ferrite filled rectangular guide.

transverse cutoff. Therefore, by demonstrating the existence of the lowest order mode we show that an infinite number of these anomalous modes may propagate simultaneously. These modes are, however, bound very tightly as surface waves to the side walls of the guide because of their strong transverse cutoff. The medium is therefore used in a very inefficient manner and high loss results, the loss increasing with mode number.

The higher propagating rectangular waveguide modes have an analogue in the higher propagating modes in a ferrite filled circular waveguide. This analogue occurs in terms of the integral number of peripheral variations. We find, similarly, an infinite number of such propagating modes each one corresponding to a given polarization sense and having a given number of peripheral variations. The reservations on practical transmission still hold in the same manner as in the rectangular case.

In the course of preparing this publication it was brought to the author's attention that Mikaelyan⁴ employed an analysis similar, in part, to that developed here. It is felt, in the present analysis, that the physical results are made more readily evident by a consideration of the limiting case of small guides, with large ratios of width to height in the case of rectangular waveguides. The choice of such large ratios is made to simplify analyses involving imaginary values of k_{x1} and k_{x2} , wherein the wave is considered to be bound to one wall of the guide and reflections from the opposite wall are of negligible amplitudes.

II. ANALYSIS OF TRANSVERSELY MAGNETIZED FERRITE IN RECTANGULAR GUIDE

The character of the ferrite medium is introduced through the Polder permeability tensor:

$$T = \begin{pmatrix} \mu & i\kappa & 0 \\ -i\kappa & \mu & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad (1)$$

The quantities μ and κ relate to the self and inductive permeabilities transverse to the z axis. The relative permeability along the z axis is given as unity. These permeabilities may be expressed as follows in gaussian units.¹

$$\mu = 1 + \frac{4\pi M_s \gamma \omega_0}{\omega_0^2 - \omega^2} \quad (2a)$$

$$\kappa = \frac{4\pi M_s \gamma \omega}{\omega_0^2 - \omega^2} \quad (2b)$$

$$\gamma = 2.80 \text{ Mc/sec/oersted}$$

$$\omega_0 = \gamma H_0$$

$$H_0 = \text{Internal dc magnetic field}$$

$$4\pi M_s = \text{Saturation magnetization}$$

Maxwell's equations are given as:

$$\text{Curl } \mathbf{H} = i\omega\epsilon\mathbf{E} \quad (3a)$$

$$\text{Curl } \mathbf{E} = -i\omega\mu_0\mathbf{T} \cdot \mathbf{H} \quad (3b)$$

Assuming a plane wave of dependence $\epsilon^{i(\omega t - \mathbf{k} \cdot \mathbf{R})}$, and appropriately combining (3a) and (3b), we have,

$$[\mathbf{k}\mathbf{k} - k^2\mathbf{I} + \omega^2\epsilon\mu_0\mathbf{T}] \cdot \mathbf{H} = 0 \quad (4)$$

The operator in square brackets is a dyadic which may be represented in matrix form. The quantity $\mathbf{I}$ is the idemfactor, having a unit diagonal representation. If we are to require that a non-trivial field $\mathbf{H}$ exist, the determinant of the operator in (4) must vanish. Since all rays traveling perpendicularly to the magnetizing axis are equivalent the medium is degenerate in the transverse plane, and some simplification is achieved in causing $\mathbf{k}$ to lie in the yz plane and letting $k_x = 0$. Some further simplification is achieved in normalizing the Polder tensor such that

$$\mathbf{T} = \frac{1}{\omega^2\epsilon\mu_0} \begin{pmatrix} f & ig & 0 \\ -ig & f & 0 \\ 0 & 0 & h \end{pmatrix} \quad (5)$$

The following secular equation is then formed.

$$\begin{vmatrix} -k^2 + f & ig & 0 \\ -ig & -k_z^2 + f & k_y k_z \\ 0 & k_y k_z & -k_y^2 + h \end{vmatrix} = 0 \quad (6)$$

Introducing the substitution $p \equiv k_z^2/k^2$, and recognizing that

$$k_y^2 = k_z^2 \left(\frac{1-p}{p} \right)$$

we have upon expanding (6),

$$p^2[(f^2 - g^2)h + k_z^2(f^2 - fh - g^2)] + p[(h - f)k_z^4 - k_z^2(f^2 + fh - g^2)] + k_z^4 f = 0. \quad (7)$$

We note, in general, two solutions in p corresponding to each value of k_z , indicating birefringence of the medium. In particular k_z must be non-vanishing for birefringence to occur or, stated alternatively, rf field gradients must exist parallel to the applied magnetic field to obtain birefringence.

The characteristic vector solutions of (4) may be expressed for each solution of (7); they are the magnetic fields,

$$\mathbf{H} = H_x \begin{pmatrix} 1 \\ i \frac{h}{g} \left(f - \frac{k_z^2}{p} \right) \\ \mp i \frac{k_z^2}{g} \left(\frac{1-p}{p} \right)^{\frac{1}{2}} \\ \frac{\left(f - \frac{k_z^2}{p} \right)}{\left(h - k_z^2 \left(\frac{1-p}{p} \right) \right)} \end{pmatrix} \epsilon^{-i(k_y y + k_z z)} \quad (8)$$

and the corresponding E fields,

$$\mathbf{E} = \frac{k_z}{\omega \epsilon} H_x \begin{pmatrix} i \frac{h}{g} \frac{\left(f - \frac{k_z^2}{p} \right)}{h - k_z^2 \left(\frac{1-p}{p} \right)} \\ -1 \\ \pm \left(\frac{1-p}{p} \right)^{\frac{1}{2}} \end{pmatrix} \epsilon^{-i(k_y y + k_z z)} \quad (9)$$

The sign indeterminacy above is defined with respect to the ratio k_y/k_z , the upper sign being given by the positive value of this ratio.

We shall analyze the rectangular waveguide by first seeking parallel plane solutions and then utilizing these solutions to form those of the rectangular guide. We choose as parallel planes those perpendicular to the applied magnetic field, or z direction and having a separation b . Because of the absolute uniformity of this type of structure, the field configurations as a function of the coordinates transverse to the magnetic field, x and y , may change only by a uniform phase factor. Again, the choice of transverse axes is made such that these phase variations occur only along y .

From (7) we would find that a specification of k_y leads to a quadratic equation in p , with an appropriate consequent multiplicity in k_z^2 . Let us define as a partial wave any standing wave in the z direction corre-

sponding to some linear combination of the positive and negative values of k_z for one of the values of k_z^2 . Examination of (9) reveals that the ratio of E_y to E_x , the field components tangent to the bounding walls, to be independent of the sign of k_z . Hence, each partial wave has an individual value of this ratio irrespective of its standing wave distribution in the z direction. It is thus impossible, in general, to provide a mutual cancellation of two or more partial waves at the electric walls by combinations of such partial waves, with the consequence that each partial wave must individually satisfy the boundary requirement. We find, then, that each partial wave takes on the familiar condition $k_z = m\pi/b$.

The parallel plane waves now will be appropriately oriented and superposed to satisfy the side wall boundary conditions in the rectangular guide. Since, as shown in Fig. 2, mutual cancellation is required on the side walls of the rectangular guide, the rate of vertical variation must be identical for all the component parallel plane waves; thus m is a constant of the waveguide mode and k_z is uniquely specified.

Two essential characteristics thus define a rectangular waveguide mode in a transversely magnetized, ferrite filled, medium.

1. The modes are ordered by integral values of m in the relationship $k_z = m\pi/b$.
2. The propagation constant k_y is uniquely specified.

Standing waves may now be formed in the z direction satisfying electric boundary conditions at the parallel planes. Each partial wave of the electric field may then be expressed as follows corresponding to its appropriate value of p :

$$E = \frac{m\pi}{b\omega\epsilon} H_x \begin{pmatrix} \frac{h \left(f - \frac{1}{p} \left(\frac{m\pi}{b} \right)^2 \right) \sin \frac{m\pi}{b} z}{g \left[h - \left(\frac{m\pi}{b} \right)^2 \left(\frac{1-p}{p} \right) \right]} \\ i \sin \frac{m\pi z}{b} \\ \left(\frac{1-p}{p} \right)^{\frac{1}{2}} \cos \frac{m\pi}{b} z \end{pmatrix} e^{-i(m\pi/b)(1-p/p)^{\frac{1}{2}}y} \quad (10)$$

Let us now specialize our analysis to the small guide case. The requirement of birefringence to produce small guide propagation demands that k_z be non-vanishing and that m take on an integral value of unity or greater. We have, from (7), the two limiting values of p corresponding

to a small value of b ,

$$p_1 = \frac{f}{f-h} = \frac{\mu}{\mu-1} \quad (11a)$$

$$p_2 = k_z^2 \frac{f-h}{f^2-fh-g^2} = \frac{k_z^2}{\omega^2 \mu_0 \epsilon} \frac{\mu-1}{\mu^2-\mu-\kappa^2} \quad (11b)$$

Discarding the z dependence in equation (10) and dropping a constant multiplier, the two characteristic electric field solutions become:

$$\mathbf{E}^{(1)} = \begin{pmatrix} \frac{1-\mu}{\kappa} \\ i \\ i\mu^{-\frac{1}{2}} \end{pmatrix} \mathcal{E}^{(m\pi/b)\mu^{-\frac{1}{2}}y} \quad (12)$$

$$\mathbf{E}^{(2)} = \begin{pmatrix} 0 \\ i \\ i \end{pmatrix} \mathcal{E}^{(m\pi/b)y} \quad (13)$$

Equations (12) and (13) are parallel plane solutions obtained for some arbitrary direction, y , transverse to the magnetic field. This direction need not be intrinsically real; mathematically, it simply satisfies Maxwell's equations. We may transform to a desired waveguide frame of reference by rotations φ_1 and φ_2 , corresponding to p_1 and p_2 , about the z axis, where these rotations may possibly be made through complex angles. We then have for the electric fields in the new space:

$$\mathbf{E}^{(1)} \rightarrow \begin{pmatrix} \frac{1-\mu}{\kappa} \cos \varphi_1 + i \sin \varphi_1 \\ -\left(\frac{1-\mu}{\kappa}\right) \sin \varphi_1 + i \cos \varphi_1 \\ i\mu^{-\frac{1}{2}} \end{pmatrix} \mathcal{E}^{(m\pi/b)\mu^{-\frac{1}{2}}(y\cos\varphi_1+x\sin\varphi_1)} \quad (14)$$

$$\mathbf{E}^{(2)} \rightarrow \begin{pmatrix} i \sin \varphi_2 \\ i \cos \varphi_2 \\ i \end{pmatrix} \mathcal{E}^{(m\pi/b)(y\cos\varphi_2+x\sin\varphi_2)} \quad (15)$$

The new y axis of the transformed coordinates is now considered the longitudinal axis of the waveguide.

The partial wave fields of (14) and (15) may be joined to form a single

mode by equating the propagation constant. Therefore,

$$\cos \varphi_2 = \mu^{-\frac{1}{2}} \cos \varphi_1 \quad (16)$$

where $\cos \varphi_2$ is imaginary for propagation. Propagation may therefore occur for $\mu > 0$ and $\cos \varphi_1$ imaginary and/or, $\mu < 0$ and $\cos \varphi_1$ real.

Boundary conditions require E_y and E_z to vanish at both guide side walls. Four equations result which may be satisfied, in turn, by a superposition of four transverse waves involving k_{x_1} , $-k_{x_1}$, k_{x_2} , and $-k_{x_2}$ corresponding to values $\pm \varphi_{1,2}$. For $\mu > 0$ both of the birefringent rays have transverse decay. Since the magnitudes of $k_{x_{1,2}}$ are large in small size guide (see Introduction) boundary conditions need be satisfied for practical purposes at only a single wall. We are then left with the simplification of only two equations in two unknowns.

Setting $x = 0$ in (14) and (15) and taking equation (16) into account, we have the boundary conditions

$$A \left[-\frac{(1-\mu)}{\kappa} \sin \varphi_1 + i \cos \varphi_1 \right] + B[i\mu^{-\frac{1}{2}} \cos \varphi_1] = 0 \quad (17)$$

$$A[\mu^{-\frac{1}{2}}] + B = 0 \quad (18)$$

With the result that

$$\cot \varphi_1 = -i \frac{\mu}{\kappa} = \left(\frac{k_y}{k_{x_1}} \right) \quad (19)$$

Choosing k_y positive real, k_{x_1} is positive imaginary for κ positive and negative imaginary for κ negative. The rf field therefore hugs the right wall for $\kappa > 0$ and the left for $\kappa < 0$, or, alternatively, switches sides in the change from a forward to backward direction of propagation.

Equation (19) may be written equivalently as

$$\cos^2 \varphi_1 = \frac{\mu^2}{\mu^2 - \kappa^2} \quad (20)$$

Propagation, occurring for imaginary values of $\cos \varphi$ and $\mu > 0$, is obtained for $|\mu| < |\kappa|$.

Let us now analyze, the possibility of small guide propagation for $\mu < 0$. We find, from (16), that $\cos \varphi_1$ is real for this case. Two cases arise; the first for which $|\cos \varphi_1| < 1$ and the second for the reverse situation.

Let us first consider the case of $|\cos \varphi_1| < 1$. From (14), k_{x_1} is real whereas from (15) k_{x_2} is imaginary. Let us associate wave amplitudes

with x dependences as follows:

$$\begin{aligned} A \varepsilon^{-ik_{x1}x} \\ B \varepsilon^{ik_{x1}x} \\ C \varepsilon^{-ik_{x2}x} \\ D \varepsilon^{ik_{x2}(x-a)} \end{aligned}$$

where a is the guide width. Let us assume that k_{x2} is a sufficiently large imaginary quantity of such sign that

$$\varepsilon^{-ik_{x2}a} \ll 1$$

This assumption will be seen to be consistent with the solution. [(26b) for small size guide.]

Setting up the boundary conditions for E_y and E_z at $x = 0$, we have from (14), (15), and (16),

$$(A - B) \left(\frac{\mu - 1}{\kappa} \right) \sin \varphi_1 + i(A + B) \cos \varphi_1 + iC\mu^{-\frac{1}{2}} \cos \varphi_1 = 0 \quad (21a)$$

$$(A + B)\mu^{-\frac{1}{2}} + C = 0 \quad (21b)$$

let $r = B/A$. Combining these last two equations we have

$$\frac{1 - r}{1 + r} = \frac{\kappa}{\mu} \cot \varphi_1 \quad (22)$$

Satisfying the boundary conditions at $x = a$ produces an equation similar to (22) with the substitution

$$r \rightarrow r \varepsilon^{-i2k_{x1}a} = r \varepsilon^{-i2\lambda}$$

Thus

$$\frac{1 - r}{1 + r} = \frac{1 - r \varepsilon^{-i2\lambda}}{1 + r \varepsilon^{i2\lambda}} \quad (23)$$

Equation (23) is satisfied by the condition $\lambda = n\pi$. Since $k_{x1} = i(m\pi/b)\mu^{-\frac{1}{2}} \sin \varphi_1$, we have

$$\sin \varphi_1 = i\mu^{\frac{1}{2}} \frac{n}{m} \frac{b}{a} \quad (24)$$

The assumption that $\cos \varphi_1$ is real and less, in magnitude, than unity is realized by the condition

$$(-\mu)^{\frac{1}{2}} \frac{n}{m} \frac{b}{a} < 1 \quad (25)$$

Only in the limiting condition of a infinitely greater than b do all modes (m, n) propagate in the negative region of μ . In this particular case, $\sin \varphi_1 = 0$ and we find from (21a) and (21b) that $C = D = 0$. Thus we find a situation in which the guide boundary conditions are satisfied by but a single class rays of the two classes available.

This result is entirely comprehensible if we observe the wave number relationship obeyed by $\mathbf{k}_1$ and $\mathbf{k}_2$. Employing the definition of p which states that $k^2 = k_z^2/p$, and using (11a) and (11b), we have

$$k_{x_1}^2 + k_{y_1}^2 + k_z^2/\mu = 0 \quad (26a)$$

$$k_{x_2}^2 + k_{y_2}^2 + k_z^2 = 0 \quad (26b)$$

As stated in the Introduction, it is an entirely consistent procedure to satisfy boundary requirements with real wave numbers over the negative range of μ using the class of rays indicated for (26a) above.

More generally, (25) shows a complex relationship of the ordering of propagation modes by n and m , for finite a , for a given negative value of μ . In contrast to the $\mu > 0$ case, propagation may possibly not occur for a range of lower order integral values of m . As μ becomes increasingly large in magnitude, m must likewise take on increasingly higher values for transmission to occur.

The case of $\cos \varphi_1$ real and greater, in magnitude, than unity, leads to trivial result. Both partial waves have imaginary values of k_x , for this case, and the far wall receives essentially no coupling. Analysis simply repeats the result of (20) and we find that $|\mu| > |\kappa|$ and $\mu < 0$. If the Polder tensor components given in (2a) and (2b) are plotted (see Fig. 5). We find that this last set of inequalities form an impossible combination.

Summarizing we find that a rectangular waveguide of any dimension (and, in particular a guide of arbitrarily small dimensions), filled with a lossless transversely magnetized ferrite medium, will support an infinite number of freely propagating modes at any frequency for which $|\mu| < |\kappa|$. The character of these modes differs considerably in the two regions of $\mu < 0$ and $\mu > 0$ and somewhat different viewpoints of propagation must be taken. We shall find similar results relating to the longitudinally magnetized ferrite filled circular waveguide in the following section.

III. ANALYSIS OF LONGITUDINALLY MAGNETIZED FERRITE IN CIRCULAR GUIDE

We now proceed to a second structural geometry in which an anomalous behavior occurs attributable to the birefringence of the medium.

This is the circular guide which has been the subject of considerable analysis by Suhl and Walker. It is instructive, however, to repeat the analysis of this case, in the small guide limit, showing more pointedly its behavior from the viewpoint of combinations of the two types of waves in the medium.

The character of transmission in undersized circular waveguide is very similar to that of the undersized rectangular case. We may demonstrate the physical significance of this statement by the following argument. The excitation in a rectangular waveguide, for $|\mu| < |\kappa|$ and $\mu > 0$, is essentially that of a surface wave bound very tightly to a single wall. Considering this wall alone, which may now be extended to arbitrary dimensions but with k_z kept large, it may be wrapped upon itself either about the magnetic field as an axis or containing the magnetic field peripherally. In either event, the wrapped guide must start and terminate at the same phase, requiring a multiplicity of 2π around the circumference, and the wave must thus continue to have a large k_z value. Considering the large value of k_z and the state of excitation of the ferrite, the small circular guide may propagate.

Analysis will demonstrate that propagation also takes place in the region $\mu < 0$. The quantity k_{x1} is real and k_{x2} imaginary, see (26), leading to a case essentially similar to that of the rectangular waveguide. The analogy is appropriate to the case of b/a of finite value for which the rectangular guide requires the appearance of both refractions. We now proceed to obtain the field solutions for the circular guide.

Referring to (9) for the plane wave solution of the electric field, let us define to within a constant multiplier.

$$\mathbf{E} = \begin{pmatrix} iE_x \\ E_y \\ E_z \end{pmatrix} e^{-i(k_y y + k_z z)} \quad (27)$$

where, for the case of large k_z (9, 12, 13)

$$\begin{aligned} E_x^{(1)} &= \frac{1 - \mu}{\kappa} & E^{(2)} &= 0 \\ E_y^{(1)} &= E_z^{(2)} = -1 \\ E_z^{(1)} &= i\mu^{-\frac{1}{2}} & E_z^{(2)} &= i \\ \frac{k_{y1}}{k_z} &= i\mu^{-\frac{1}{2}} & \frac{k_{y2}}{k_z} &= i \end{aligned}$$

We shall consider here, of the two possible wrapped-wall structures, that case in which the magnetic field is applied axially as shown in Fig. 3. Referring to Fig. 4, the cylindrical drical electric wave satisfying Max-

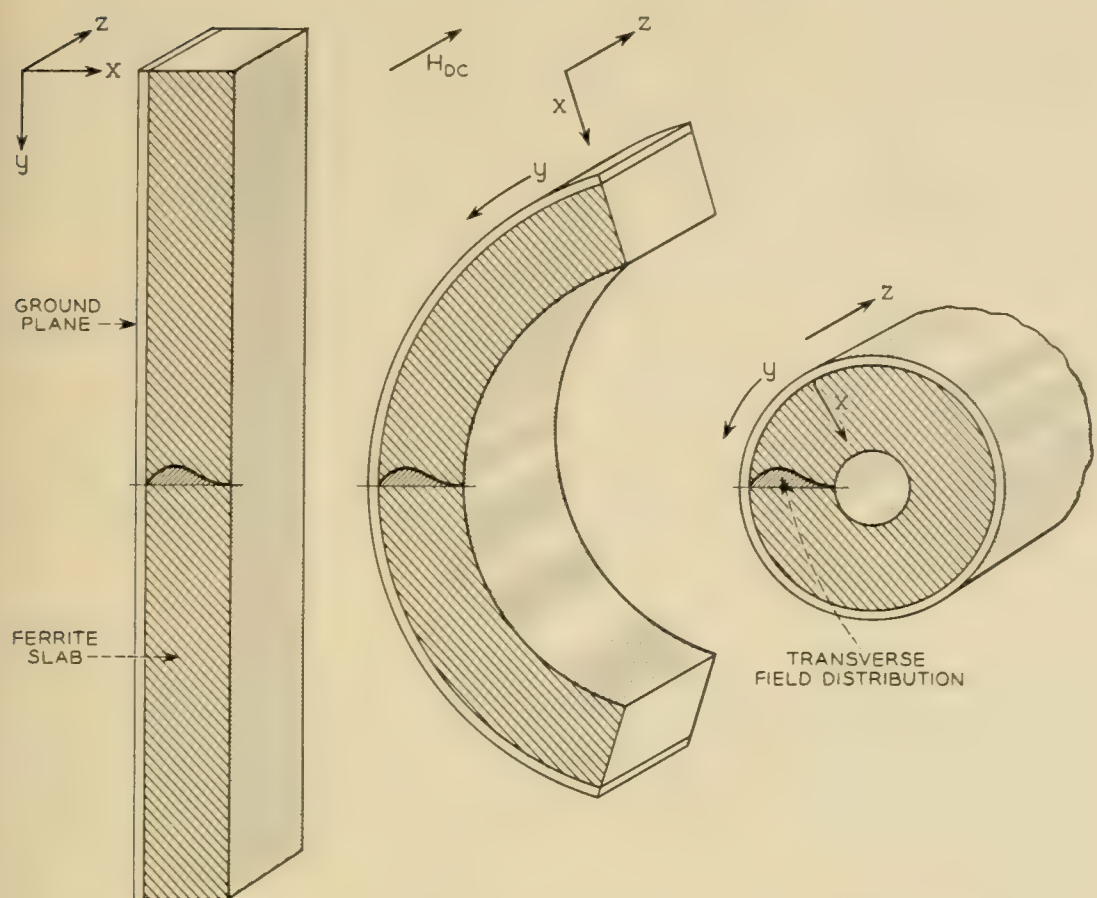


Fig. 3 — Axially magnetized filled circular guide formed by wrapping wall.

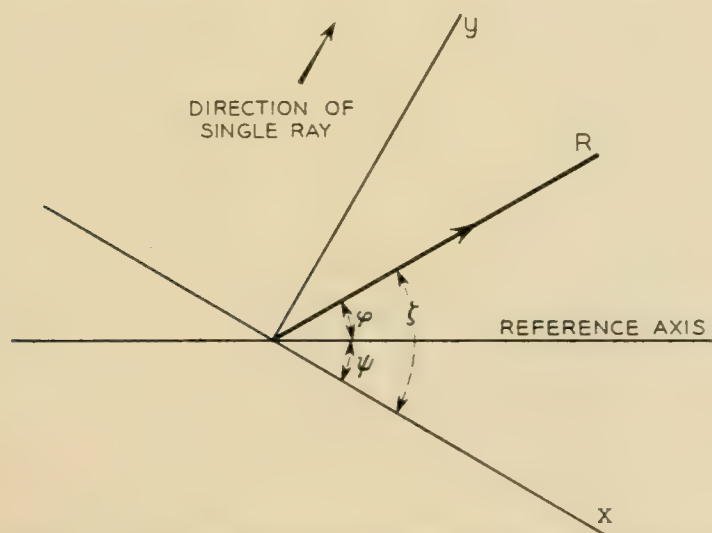


Fig. 4 — Transformation to polar coordinate frame.

well's equations and the boundary conditions for this structure is obtained by integrating plane waves of the form of (27) traveling at all possible angles ψ , the integration being subject to a weighting factor $G(\psi)$ to obtain the most general field. The coordinates (r, φ) refer to the physical system and the coordinate ψ identifies a plane wave traveling along a particular y axis. We have thus in an (r, φ, z) coordinate frame:

$$\underline{E} = \frac{1}{2\pi} \int_0^{2\pi} G(\psi) \begin{pmatrix} \cos \zeta & \sin \zeta & 0 \\ -\sin \zeta & \cos \zeta & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} iE_x \\ E_y \\ E_z \end{pmatrix} \varepsilon^{-i(k_y y + k_z z)} d\psi \quad (28)$$

Recognizing that

$$d\psi = d\zeta$$

$$y = r \sin \zeta$$

and

$$G(\psi) = G(\zeta - \varphi)$$

an integration results over the variable ζ . Because of the uniqueness of the field as a function of φ , the only term containing φ , $G(\zeta - \varphi)$, must be a periodic function in its argument. A typical mode is formed by choosing one of the terms of the Fourier series of $G(\zeta - \varphi)$, namely $\varepsilon^{in(\zeta - \varphi)}$.

We find from (28) that

$$\underline{E}_n = \begin{pmatrix} i \left[E_x \frac{n}{\rho} J_n(\rho) + E_y J_n'(\rho) \right] \\ E_x J_n'(\rho) + E_y \frac{n}{\rho} J_n(\rho) \\ E_z J_n(\rho) \end{pmatrix} \varepsilon^{-i(k_z z + n\varphi)} \quad (29)$$

where $\underline{E}_n$ is that partial expansion of the total field $\underline{E}$, corresponding to the number of angular variation n , and $\rho = k_y r$. There are two values of ρ corresponding to the two values of k_y , and each leads to a partial wave. Let A and B be the respective partial wave amplitude; satisfying the boundary conditions on E_φ and E_z , we have from (29):

$$A \left(\frac{1 - \mu}{\kappa} \right) J_n'(\rho_1) - \frac{nA}{\rho_1} J_n(\rho_1) - J_n(\rho_2) = 0 \quad (30a)$$

$$A \mu^{-\frac{1}{2}} J_n(\rho_1) + B J_n(\rho_2) = 0 \quad (30b)$$

where ρ_1 and ρ_2 are defined for $r = R$, the radius of the cylinder. Recognizing

nizing that $\rho_1 = \mu^{-\frac{1}{2}}\rho_2$, we have from (30)

$$\frac{\mu}{\kappa} J_n'(\rho_1) + n \frac{J_n(\rho_1)}{\rho_1} = 0 \quad (31)$$

where $\rho_1 = ik_z \mu^{-\frac{1}{2}} R$.

Equation (21) may be modified by a recurrence relationship to become

$$\frac{\kappa}{\mu} + 1 = \frac{\rho_1}{n} \frac{J_{n+1}(\rho)}{J_n(\rho_1)} \quad (32)$$

For $\mu > 0$ the quantity ρ_1 is a pure imaginary for large real values of k_z . Since the n^{th} order Bessel function is monotonic in imaginary arguments and possesses the multiplier $(i)^n$, the right-hand side is negative for n positive. For $n > 0$, propagation occurs for

$$|\kappa| > |\mu|$$

$$\text{sgn } \kappa = -\text{sgn } \mu$$

Inspection of (31) reveals that a reversal of the sign of n is equivalent to reversing the sign of κ . This conforms to the physical situation in which reversal of the sense of circular polarization is equivalent to the reversal of magnetic field. Thus for $n < 0$ and $\mu > 0$,

$$|\kappa| > |\mu|$$

$$\text{sgn } \kappa = \text{sgn } \mu$$

We find, from the above arguments, that just one sense of circular polarization propagates in an undersized circular guide for $\mu > 0$ and for a given direction of the magnetic field. It will be demonstrated shortly that propagation occurs for $\mu < 0$, but with an entirely different structure of modes. The right-hand side of (32) is monotonic as a function of ρ for $\mu > 0$, leading to only one solution for each value of n . This will not be the case for $\mu < 0$.

It is of interest first, however, to observe the limiting approach to $\mu = 0$ in the region of $\mu > 0$. The right-hand side of (32) is finite for finite imaginary values of ρ_1 , so that the only solution as μ approaches zero is that for which the magnitude of ρ_1 becomes infinitely great. The Bessel function is asymptotically expandible as a cosine divided by a square root of its argument. Thus

$$J_n(\rho_1) = \frac{1}{2} \sqrt{\frac{2\pi}{\rho_1}} (\varepsilon^{i(\rho_1 + [2n+1](\pi/4))} + \varepsilon^{-i(\rho_1 + [2n+1](\pi/4))}) \quad (33)$$

* Equation (31) may likewise be obtained from the small radius limit in (34) of Reference 2.

Considering ρ_1 to be positive imaginary, as $\rho_1 \rightarrow i\infty$, (32) becomes

$$\frac{\kappa}{\mu} = \frac{-i\rho_1}{n} \quad (34)$$

Substituting for ρ_1 , we have

$$k_z = \frac{n\kappa}{\mu^{\frac{1}{2}}R} \quad (35)$$

Thus, as μ approaches zero from values greater than zero, the propagation constant tends to become singular. Physically, however, μ does not vanish but approaches a small imaginary value caused by ferrite losses. The propagation constant k_z becomes complex and takes on a large imaginary component, signifying large guide attenuation. Since these losses occur in the limited neighborhood of $\mu = 0$, we may construe this waveguide behavior as corresponding to a system resonance.

In the region $\mu < 0$, ρ_1 becomes real while ρ_2 remains imaginary. The right-hand side of (32) is now composed of only real arguments. Since the zeros of different order Bessel functions alternate, the right side of (32) contains a succession of poles and zeros, leading to an infinite number of branches with each containing a solution ρ_1 to the equation. Thus there are an infinite number of propagating modes corresponding to each value

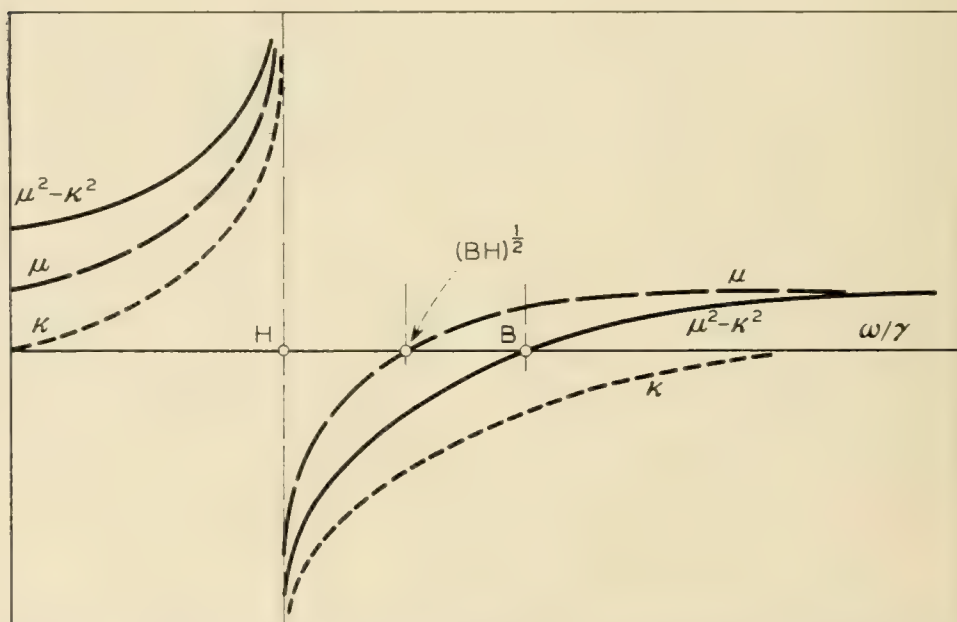


Fig. 5 — Frequency characteristics of Polder tensor components.

of n , in marked contrast to the case of $\mu > 0$. The solutions remain identical, as before, if both κ and n are simultaneously reversed in sign, but differ if only one of the two quantities is reversed.

Since $\mu = 0$ is a branch point, the limiting condition as μ approaches zero for values $\mu < 0$ differs from that for the reverse case. Equation (32) is now satisfied in the limit of small μ by the real zeros of $J_n(\rho_1)$. Since these roots are finite, k_z , equal to $-(-\mu)^{\frac{1}{2}}\rho_1/R$, tends towards zero for all modes. Since the formulae developed in this paper always presume large wave numbers, we may infer a vanishing value of k_z to simply represent a value which is small relative to the reciprocal of the waveguide radius. In any event, k_z is no longer singular at $\mu = 0$, and there is no resonance in the approach from negative values of μ .

In sum, the features of the circular guide strongly resemble those of the rectangular guide in the region of $\mu > 0$. This was to be anticipated by the "wrapped wall" construction where the wave is tightly bound to the wall. The wrapped equivalences do not hold in the region $\mu < 0$ since, with harmonic transverse dependence, the wave is no longer bound to the wall. This lack of equivalence is manifested in the matter of ordering modes. For a rectangular waveguide of finite aspect ratio, we find from (25) that there are but a finite number of modes corresponding to each value of m for $\mu < 0$. The circular guide differs in providing an infinite number of modes corresponding to each value of n . Further, whereas the circular guide covers the entire range of $|\mu| < |\kappa|$, (25) indicates that the various modes of the rectangular guide covers a more restricted range determined by the guide aspect ratio.

IV. CONCLUSIONS

The waveguide behavior analyzed in this paper has been experimentally observed⁵ and good correlation has been obtained. From the viewpoint expressed of forming a guide cross-section by wrapping a wall to which a surface wave is bound, we may anticipate that the unusual behavior observed in the two types of guides examined is probably characteristic of many other structures.

It is not clear, at this time, if the complete set of modes of either the rectangular or circular guides have been exhausted. We already observe that an infinite number of modes propagate simultaneously so that scattering problems become considerably more complex than in the usual cases. It is felt by the author that the field of waveguide analysis calls for new methods and techniques of modal synthesis when ferrite loaded structures are considered.

ACKNOWLEDGMENT

I feel particularly indebted to R. C. Fletcher and H. Boyet for their many valuable comments and suggestions.

REFERENCES

1. D. Polder, *Phil. Mag.*, **40**, p. 99, Jan., 1949.
2. H. Suhl and L. R. Walker, *B.S.T.J.*, **33**, 3, May, 1954. See Fig. 9g, p. 615, and Table I, p. 642.
3. G. H. B. Thompson, *Nature*, **175**, p. 1135, June 25, 1955.
4. A. L. Mikaelyan, *Doklady, A. N. USSR*, **98**, 6, pp. 941-944, 1954.
5. H. Seidel, *Proc. I.R.E.*, **44**, p. 1410, Oct., 1956.

Measurement of Dielectric and Magnetic Properties of Ferromagnetic Materials at Microwave Frequencies

By WILHELM VON AULOCK and JOHN H. ROWEN

(Manuscript received August 15, 1956)

Some experimental techniques are discussed which permit measurement of the magnetic and dielectric properties of ferrite materials in the microwave region by observing the perturbation in a cylindrical cavity due to insertion of a small ferrite sample. A comparison of the properties of thin disc samples with those of small spheres shows that discs yield more accurate results in the region below ferromagnetic resonance whereas spheres are preferable for the study of ferrite properties near resonance. A short description of instrumentation for cavity measurements at 9,200 mc is given and experimental results of disc measurements are reported for a low-loss BTL ferrite and several disc diameters. A comparison of experimental results with Polder's theory indicates that the loss of polycrystalline ferrites below resonance is considerably lower than that predicted from an evaluation of the width of the resonance absorption line.

1. INTRODUCTION

The dielectric and magnetic properties of semi-conducting ferromagnetic materials such as ferrites have been the subject of intense study in recent years. Analytical expressions for the components μ and κ of the permeability tensor of a loss-free single-crystal ferrite were derived by Polder.¹ These expressions were later modified to include a loss factor α .^{2, 3} Yager and others⁴ measured the resonance absorption of single crystals of nickel ferrite and found very good agreement with theory provided the loss factor α was determined from the width of the measured resonance absorption line. However, when Artman and Tannenwald⁵ measured the real and imaginary parts of μ and κ for polycrystalline ferrites they found that agreement with theory was somewhat less than perfect if α was also determined from the measured line width. Discrepancies were observed for both real and imaginary parts of $\mu +$

κ in the region below resonance because the effects of polycrystalline structure and anisotropy forces were neglected in Polder's and Hogan's² analysis. Furthermore, it was assumed in the derivation of the permeability tensor that the ferrite is saturated with a biasing dc magnetic field which is large compared to the microwave magnetic field.

There exists a great need for experimental data for all those conditions where some of the above assumptions do not hold. In particular, the region of biasing magnetization between zero and ferromagnetic resonance is of interest because it is the operating region for many ferrite devices such as phase shifters, modulators, and field displacement isolators. Techniques for the measurement of ferrite parameters below resonance were investigated and it was found that the measurement of the perturbation of a degenerate cylindrical cavity by a thin ferrite disc yielded accurate results, whereas observation of the cavity perturbation caused by a small sphere produced less accurate data.

It is the purpose of this paper to describe and discuss the thin disc method and to compare it with other techniques described in the literature.^{5, 7} After defining the ferrite parameters as constants in Maxwell's equations it is shown how these parameters can be obtained from various measuring techniques. Instrumentation for the thin disc technique is described and a few remarks are made pertaining to experimental difficulties. Finally, some measurements of low-loss ferrites are reported and compared with values predicted by Polder's relations.

2. DESCRIPTION OF FERRITE PARAMETERS

It is customary⁸ to define the electric and magnetic polarization vectors $\vec{P}$ and $\vec{M}$ in terms of the field vectors $\vec{E}$ (electric field intensity), $\vec{D}$ (electric displacement), $\vec{H}$ (magnetic field intensity), and $\vec{B}$ (magnetic induction). In the M.K.S. system we have:

$$\vec{P} = \vec{D} - \epsilon_0 \vec{E}$$

$$\vec{M} = \vec{B}/\mu_0 - \vec{H}$$

$$\epsilon_0 = 8.854 \times 10^{-12} \text{ farad/meter, permittivity of free space}$$

$$\mu_0 = 4\pi \times 10^{-7} \text{ henry/meter, permeability of free space}$$

Then, the intrinsic parameters of a ferrite medium are defined as those quantities which relate $\vec{P}$ and $\vec{M}$ to the electric and magnetic fields in the medium respectively.

$$\vec{P} = \epsilon_0 \chi_e \vec{E}$$

$$\vec{M} = \vec{\chi}_m \vec{H}$$

Whereas the electric susceptibility χ_e is a scalar quantity in ferrites the

magnetic susceptibility $\vec{\chi}_m$ is known to have tensor properties. Assuming that the static magnetic field H_z is in the z -direction we have

$$\vec{\chi}_m = \begin{vmatrix} \chi_m & -j\kappa & 0 \\ j\kappa & \chi_m & 0 \\ 0 & 0 & 0 \end{vmatrix} \quad (1)$$

If we restrict ourselves to sinusoidal time variation of the RF fields we may describe electric and magnetic losses in the ferrite by regarding χ_e , χ_m , and κ as complex quantities:

$$\begin{aligned} \chi_e &= \chi_e' - j\chi_e'' \\ \chi_m &= \chi_m' - j\chi_m'' \\ \kappa &= \kappa' - j\kappa'' \end{aligned}$$

Thus, it is seen that the RF properties of a ferrite medium regardless of geometry are completely described by six "intrinsic" parameters, χ_e' , χ_e'' , χ_m' , χ_m'' , κ' , and κ'' . The dielectric constant ϵ and permeability $\vec{\mu}$ of the material are obtained from

$$\begin{aligned} \epsilon &= \chi_e + 1 \\ \vec{\mu} &= \vec{\chi}_m + 1 \end{aligned}$$

where **1** is the unit matrix.

It is the objective of the measurement to obtain each of the above parameters as a function of one or more variables of interest such as frequency, saturation magnetization, static magnetic field, temperature, applied power and others. Measurements may be made on single ferrite crystals — mostly for research purposes — or on polycrystalline material for many purposes in connection with the development of ferrite materials and devices. These measurements are generally compared to the behavior of χ_m and κ as predicted by Polder's equations. One obtains from the equation of motion of magnetization¹

$$\begin{aligned} \chi_m &= \mu - 1 = \frac{\gamma^2 M_z H_z}{\gamma^2 H_z^2 - \omega^2} \\ \kappa &= -\frac{\omega |\gamma| M_z}{\gamma^2 H_z^2 - \omega^2} \end{aligned}$$

Using Suhl and Walker's³ notation this may be written as

$$\begin{aligned} \chi_m &= \frac{p\sigma}{\sigma^2 - 1} \\ \kappa &= -\frac{p}{\sigma^2 - 1} \end{aligned}$$

where $p = |\gamma| M_z / \omega$ normalized saturation magnetization

$\sigma = |\gamma| H_z / \omega$ normalized static magnetic field in the ferrite

$|\gamma| = 2.8$ mc per oersted, gyromagnetic ratio

ω = operating frequency

It appears that some cavity techniques measure the eigenvalues of (1), $\chi_m + \kappa$ and $\chi_m - \kappa$, directly, which are seen to be

$$\chi_m \pm \kappa = \frac{p}{\sigma \pm 1} \quad (2)$$

Suhl and Walker show that a loss term may be introduced by replacing σ by $\sigma + j\alpha(\text{sgn } p)$ in (2).^{*} Separating real and imaginary parts we get

$$\chi_m' \pm \kappa' = \frac{p(\sigma \pm 1)}{(\sigma \pm 1)^2 + \alpha^2} \quad (3)$$

$$\chi_m'' \pm \kappa'' = \frac{\alpha p(\text{sgn } p)}{(\sigma \pm 1)^2 + \alpha^2} \quad (4)$$

For the determination of α from measurements it is convenient to define a loss tangent

$$\delta_{\pm} \equiv \frac{\chi_m'' \pm \kappa''}{\chi_m' \pm \kappa'} = \frac{\alpha(\text{sgn } p)}{\sigma \pm 1} \quad (5)$$

Typical curves for $\chi_m \pm \kappa$ and $\delta_{\pm}$ assuming $p = 0.5$ and $\alpha = 5 \times 10^{-2}$ are shown on Figure 1.[†] For the purpose of describing and comparing experimental results, it may be convenient to distinguish among various regions of H_z as indicated on the graph because a different measurement technique may be required for accurate measurements in each region.

3. METHODS FOR MEASURING MAGNETIC PROPERTIES

Three measurement methods have been reported in the literature all of which employ the detuning and change in $1/Q$ of a resonant cavity by a small ferrite sample. Van Trier⁷ used very thin long cylindrical samples in a coaxial cavity. Artman and Tannenwald⁵ employed small spheres, and we used thin discs⁹ both placed close to the endwall of a cylindrical degenerate cavity excited by a TE_{111} mode (Figs. 2 and 3). Recently, Berk and Lengyel⁶ suggested the use of a cylindrical post at the center

^{*} By definition $\text{sgn } p = +1$ for $p > 0$ and $\text{sgn } p = -1$ for $p < 0$.

[†] Since it is customary to use $\chi_m + \kappa$ for the designation of the resonance line this notation has been used here. Consequently, p and σ should be assumed negative.

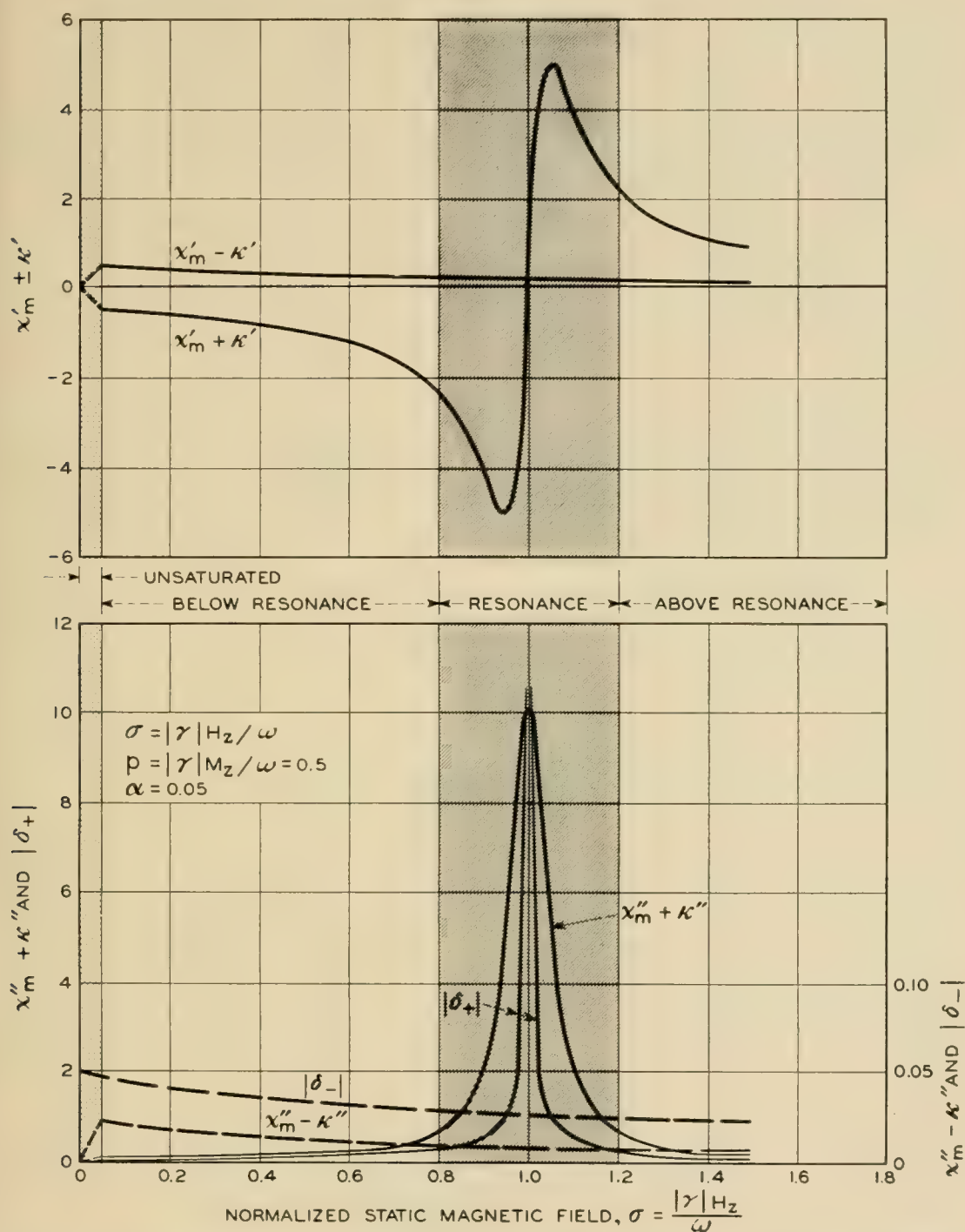


Fig. 1 — Theoretical values of $\chi_m \pm \kappa$ and loss tangent $\delta_{\pm}$ versus normalized static field σ .

of a degenerate rectangular cavity.* In principle, all these methods permit the determination of the six parameters χ_e' , χ_e'' , χ_m' , χ_m'' , κ' , and κ'' . However, in practical applications there are significant differences, e.g.,

* As this paper was being written another variation of the thin cylinder technique using the TM_{110} mode in a circular cavity was reported by Spencer and LeCraw at the I.R.E. Convention, New York, Mar. 21, 1956.

between the use of a spherical sample and a thin disc, such as the perturbing effect on the cavity field, the accuracy of small loss measurement, the occurrence of resonance at a static field where the intrinsic parameters are not at resonance,¹⁰ and the possibility of making accurate measurements of the electric susceptibility. Therefore, the sphere method and the disc method will be reviewed briefly and an attempt will be made to compare their capabilities. It is hoped that this comparison will also be helpful for the evaluation of other measuring techniques not covered in this paper.

3.1 The Small Sphere Method

In order to appreciate the significance of the quantities measured with the small sphere method it is expedient to define an effective susceptibility tensor $\vec{\chi}_{ms}$ by relating the magnetization vector* to the applied

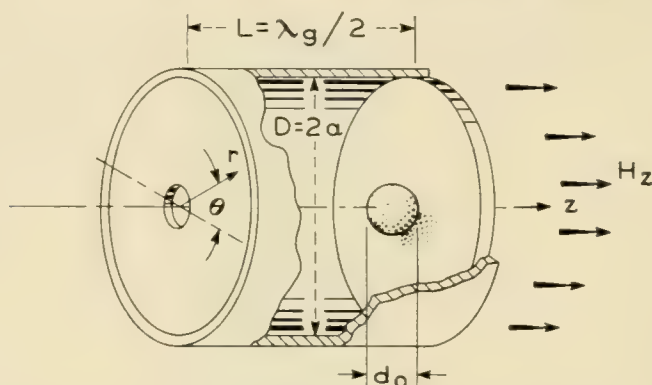


Fig. 2 — Degenerate TE_{111} cylindrical cavity with ferrite sphere.

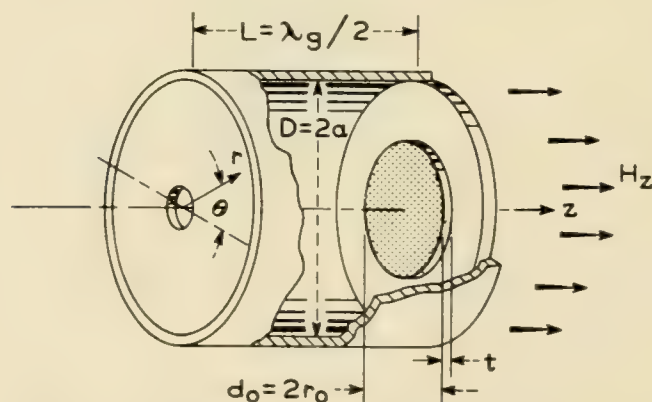


Fig. 3 — Degenerate TE_{111} cylindrical cavity with ferrite disc.

* R. A. Waldron (Institute of Electrical Engineers, Convention Oct. 29 to Nov. 2, on Ferrites, London, 1956) related the magnetic induction vector $\vec{B}$ in a

field $\vec{H}^0$

$$\vec{M} = \vec{\chi}_{ms} \vec{H}^0$$

and observing that the RF components of $\vec{M}$ and $\vec{H}^0$ (denoted by lower case letters) are related in a cylindrical coordinate system by

$$\begin{aligned} m_r &= \chi_{ms} h_r^0 - j \kappa_s h_\theta^0 \\ m_\theta &= j \kappa_s h_r^0 + \chi_{ms} h_\theta^0 \end{aligned} \quad (6)$$

Placing the small ferrite sphere close to the endwall of the cavity (Fig. 2) and observing the splitting of the resonance into two frequencies $\omega_\pm$ (related to the positive and negative circularly polarized modes) and the two changes in $1/Q$ of the cavity after application of a static magnetic field in the axial direction leads to two measurable quantities

$$\Delta\omega_\pm = \omega_0 - \omega_\pm \quad \text{frequency shift}$$

$$\Delta(1/Q)_\pm = 1/Q_\pm - 1/Q_0 \quad \text{change in internal } Q \text{ of the cavity}$$

where ω_0 and Q_0 are resonance frequency and Q of the empty cavity. It can be shown that real and imaginary part of χ_{ms} and κ_s can be obtained from†

$$\frac{2\Delta\omega_\pm}{\omega_0} = 0.6982 \frac{\lambda_0^2}{D^2} \cdot \frac{d_0^3}{L^3} (\chi_{ms}' \pm \kappa_s') \quad (7)$$

$$\Delta(1/Q)_\pm = 0.6982 \frac{\lambda_0^2}{D^2} \cdot \frac{d_0^3}{L^3} (\chi_{ms}'' \pm \kappa_s'') \quad (8)$$

The quantities d_0 , D , and L are the sphere diameter, cavity diameter and length respectively, λ_0 is the wavelength in free space associated with ω_0 . In order to obtain the intrinsic parameters χ_m and κ from (7) and (8) one may use the relationships⁶

$$\vec{H}^0 = \vec{H} + \vec{M}/3 \quad (9)$$

$$\chi_{ms} \pm \kappa_s = \frac{3(\chi_m \pm \kappa)}{\chi_m \pm \kappa + 3} \quad (10)$$

sphere to the applied field by writing $\frac{1}{\mu_0} \vec{B} = \vec{\mu}_s \vec{H}^0$ where $\vec{\mu}_s$ may be designated the external relative permeability tensor. It can be readily shown that Waldron's results are in agreement with ours if one notes that $(2/3)\vec{\chi}_{ms} = \vec{\mu}_s - \mathbf{1}$. We found that the use of the effective susceptibility tensor $\vec{\chi}_{ms}$ is much to be preferred over $\vec{\mu}_s$ because it simplifies notation and interpretation of experimental results in terms of the intrinsic quantities χ_m and κ .

† Equations (7) and (8) are identical to Artman and Tannenwald's⁵ expressions, if $4\pi^2 D^2 / (13.56 + \pi^2 D^2 / L^2)$ is substituted for λ_0^2 .

Equation (10) has a pole at $\chi_m \pm \kappa = -3$ which simply indicates the resonance condition for a sphere as derived by Kittel.¹⁰ This can be verified from (2):

$$\chi_m \pm \kappa = \frac{p}{\sigma \pm 1} \quad (11)$$

Note that p and σ are either both positive or both negative, hence, only one of the two quantities $(\chi_m + \kappa)$ or $(\chi_m - \kappa)$ goes through resonance at $|\sigma| = 1$. A similar situation exists for $\chi_{ms} \pm \kappa_s$ expressed in terms of (11)

$$\chi_{ms} \pm \kappa_s = \frac{3p}{p + 3(\sigma \pm 1)} \quad (12)$$

One of the two quantities $(\chi_{ms} \pm \kappa_s)$ goes through resonance at

$$|\sigma|_s = 1 - |p|/3 \quad (13)$$

Observing that the field in the sphere is given by (9) this may be written as

$$\frac{\omega_r}{|\gamma|} = H_z + M_z/3 = H_z^0 \quad (14)$$

(Kittel's resonance frequency of a ferrite sphere)

It is easily seen that this resonance of the spherical sample makes the evaluation of χ_m and κ from (10) rather unattractive because one would expect inaccurate results for χ_m and κ in the vicinity of the sphere resonance σ_s . Furthermore, for all numerical computations (10) must be separated into real and imaginary parts

$$\chi_{ms}' \pm \kappa_s' = 3 \frac{(\chi_m' \pm \kappa' + 3)(\chi_m' \pm \kappa') + (\chi_m'' + \kappa'')^2}{(\chi_m' \pm \kappa' + 3)^2 + (\chi_m'' \pm \kappa'')^2} \quad (15)$$

$$\chi_{ms}'' \pm \kappa_s'' = 9 \frac{\chi_m'' \pm \kappa''}{(\chi_m' \pm \kappa' + 3)^2 + (\chi_m'' \pm \kappa'')^2} \quad (16)$$

In general higher order terms of χ_m'' and κ'' may be neglected, but in the vicinity of sphere resonance these terms predominate as the term $\chi_m' \pm \kappa' + 3$ vanishes. It can be seen from the preceding discussion that the determination of χ_m' , χ_m'' , κ' , and κ'' from χ_{ms} and κ_s has its difficulties. Fortunately, there is an easier way to the interpretation of χ_{ms} and κ_s in terms of the intrinsic parameters χ_m and κ . We use (9) to define a new quantity σ' in terms of the applied magnetic field.

$$\sigma' = \sigma + p/3 = |\gamma| H_z^0 / \omega \quad (17)$$

Now (12) may be written in terms of the normalized applied static field as

$$\chi_{ms} \pm \kappa_s = \frac{p}{\sigma' \pm 1} \quad (18)$$

Comparison with (11) shows that the functional dependency of the effective parameters χ_{ms} and κ_s on the applied field H_z^0 is identical to Polder's relations for $\chi_m \pm \kappa$. This permits plotting $\chi_{ms} \pm \kappa_s$ versus H_z^0 and relabeling the coordinates $\chi_m \pm \kappa$ and H_z , provided that Polder's equations hold exactly. It would appear however that one of the reasons for measuring ferrite parameters is the fact that Polder's equations are known to be approximations which do not always hold and which are subject to many restrictions as pointed out earlier in this paper. Summing up, one may say that measurements of the parameters of a small ferrite sphere lead to excellent results in terms of effective quantities χ_{ms} and κ_s as a function of the applied static field, but may only be approximations when interpreted in terms of the intrinsic parameters χ_m and κ .

3.2 The Thin Disc Method

The use of a thin disc rather than a small sphere as a cavity perturbation eliminates most of the difficulties enumerated in the previous section because the intrinsic parameters $\chi_m \pm \kappa$ are measured directly. Placing a thin disc against the endwall of a cylindrical TE_{111} mode cavity (Fig. 3) and observing the splitting of the resonance frequency and change in $1/Q$ as before one obtains the following relationships (see Appendix for derivation)

$$\frac{\Delta\omega_{\pm}}{\omega_0} = \frac{1}{4} \frac{\lambda_0^2 t}{L^3} (\chi_m' R_1 \pm \kappa' R_2) \quad (19)$$

$$\Delta \left(\frac{1}{Q} \right)_{\pm} = \frac{1}{2} \frac{\lambda_0^2 t}{L^3} (\chi_m'' R_1 \pm \kappa'' R_2) \quad (20)$$

The quantity t denotes the thickness of the disc, and R_1 and R_2 are functions of the geometry which take into account that the RF magnetic field is not constant over the face of the disc. A plot of R_1 and R_2 versus the ratio of disc diameter to cavity diameter (Fig. 4) shows that the functions are closely equal for disc diameters less than $\frac{1}{2}D$. This implies circular polarization of the magnetic field in this region whereas the field becomes elliptically polarized as one approaches the outer diameter of the cavity. It might be argued that the disc should be small enough

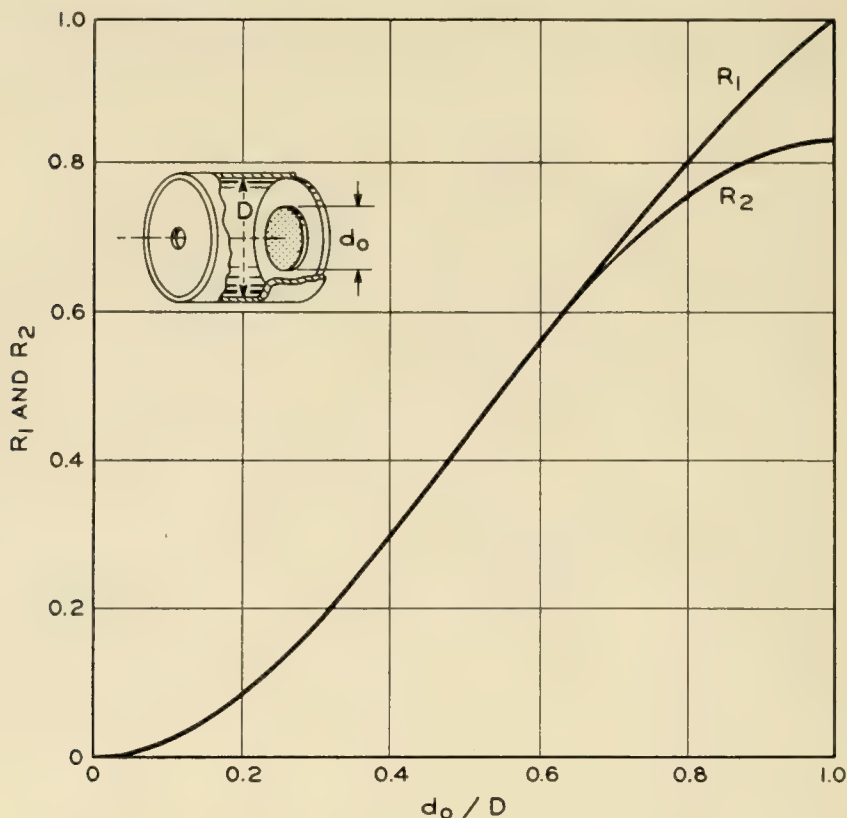


Fig. 4 — The functions R_1 and R_2 versus the ratio of disc diameter d_0 to cavity diameter D .

to be entirely within the circularly polarized region. However, this requirement appears to assume that there is some interaction, such as spinwave coupling, between adjacent regions in the disc. It is possible that such interaction exists in single crystals and leads to multiple resonance effects. However in polycrystalline material it is safe to assume that each crystallite reacts independently with the RF magnetic field, enabling us to sum these effects by simple integration.

This integration has been performed in the derivation of (19) and (20) under the assumption that the disc is thin enough to introduce only a first order perturbation into the cavity field. It will be shown presently that these assumptions are consistent with experimental results.

The measured effects $\Delta\omega_{\pm}$ and $\Delta(1/Q)_{\pm}$ depend on the volume of the perturbing body, hence one would expect that these effects are at least an order of magnitude larger for thin discs than for small spheres. This permits very accurate measurement of ferrite parameters in the region below saturation and below and above resonance. In particular, the loss parameters $\chi_m'' \pm \kappa''$ of modern low-loss ferrites can be determined accurately in these regions. However, in the resonance region, the measured effects become so large that measurement becomes difficult. Hence,

we do not recommend this method for measurement of resonance linewidth. The greatest advantage of the thin disc method lies in its ability to explore the region below resonance where many ferrite devices operate. It will be noted that this region is wider for the disc than for a small sphere or for a thin cylinder magnetized normal to its axis. Using Kittel's relations we find the field in the ferrite at which resonance occurs for a given saturation magnetization and operating frequency

$$\begin{array}{lll} H_{\text{res}} = \frac{\omega}{|\gamma|} & H_{\text{res}} = \frac{\omega}{|\gamma|} - \frac{M_z}{3} & H_{\text{res}} = \frac{\omega}{|\gamma|} - \frac{M_z}{2} \quad (21) \\ \text{Thin Disc} & \text{Small Sphere} & \text{Thin Cylinder} \end{array}$$

It is seen that for some materials a small sphere or a thin cylinder may be at resonance even before it is saturated.

4. METHODS FOR MEASURING THE DIELECTRIC PROPERTIES

In principle the electric susceptibility χ_e may be measured with any of the three sample shapes discussed above provided the sample is placed into a region of maximum electric and vanishing magnetic field, e.g., the center of a TE_{111} mode cylindrical cavity. A simple computation shows that a small sphere located at the center of the cavity produces a frequency shift and change in $1/Q$ as follows:

$$\frac{\Delta\omega}{\omega_0} = 4.189 \frac{\chi_e'}{\chi_e' + 3} \cdot \frac{d_0^3}{D^2L} \quad (22)$$

$$\Delta(1/Q) = 25.136 \frac{\chi_e''}{(\chi_e' + 3)^2} \cdot \frac{d_0^3}{D^2L} \quad (23)$$

In actual measurements it is found that the dielectric constant is not entirely independent of the diameter of the sphere, and that the measurement of the loss factor χ_e'' is difficult because the change in $1/Q$ is very small. An additional inherent difficulty of this method is the fact that (22) is rather insensitive to large values of χ_e as are frequently encountered with ferrites. Again the thin disc method permits greater ease of measurement, larger effects, and a simpler relationship between the observed quantities and χ_e . We find (see Appendix A)

$$\frac{\Delta\omega}{\omega_0} = \chi_e' \frac{t}{L} R_1 \quad (24)$$

$$\Delta(1/Q) = 2\chi_e'' \frac{t}{L} R_1 \quad (25)$$

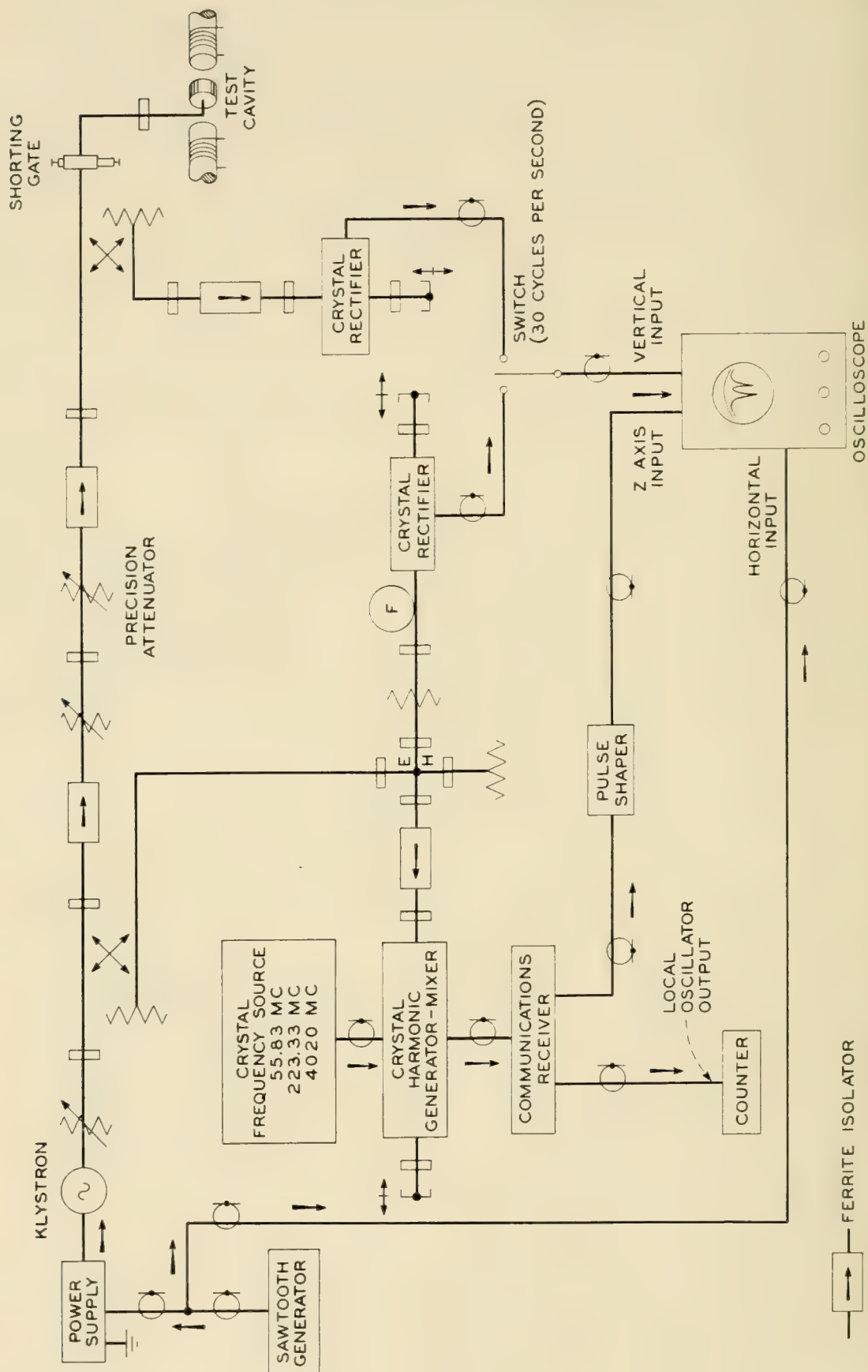


Fig. 5 — Microwave circuit for measuring ferrite parameters.

5. INSTRUMENTATION

Our measurement technique has been influenced by a number of practical considerations including the need for determining accurately and quickly a figure of merit for a large number of different ferrite materials. In particular, a variety of low loss materials has become available in experimental quantities requiring a precise technique for measuring small loss factors below resonance as a guide for further ferrite development. Therefore we were faced with the problem to develop an instrumentation capable of measuring these small loss factors but simple enough to be operated without detailed knowledge of microwave techniques. Fortunately the use of thin discs permits us to introduce a fairly large volume of ferrite into the cavity without violating the basic assumption of a small perturbation. As a consequence frequency shifts of the order of 10 mc are obtained at static magnetic fields just sufficient to saturate the material. Thus the quantities χ_m' and κ' may be measured without difficulty in the regions below and above resonance.

The thickness of the disc should be chosen to attain an aspect ratio (diameter/thickness) of 50 or larger. Discs of 0.005 to 0.007-inch thickness were employed in actual measurements at 9200 mc. For a typical measurement of a 0.005-inch disc at 9200 mc, a change of $1/Q$ of about 2 per cent corresponds to a loss term $\chi_m'' + \kappa'' = 2 \times 10^{-3}$. Measurement of Q with a reproducibility of 1 per cent has been accomplished initially by careful work, and it has been our objective to maintain this accuracy in routine measurements by semiskilled operators. This required the use of rather elaborate circuitry for the precise measurement of the changes in $1/Q$ of the cavity. Although most of these techniques have been used before in the field of microwave spectroscopy we hope that the description of this instrumentation will be of interest.

Fig. 5 shows a block diagram of the circuit. A klystron Type V58 (Varian Associates) is swept through a frequency band of about 80 mc at X-band. The resulting signal with a center frequency of 9,200 mc is used to excite a TE_{111} mode cylindrical cavity. Incident and reflected signals are separated by means of directional couplers and displayed on an oscilloscope. Both signals can be aligned with the aid of a shorting gate and a precision attenuator in front of the cavity. The reflected signal shows clearly the cavity resonance which splits into two if a ferrite disc is placed against the endwall of the cavity and magnetized along the cavity axis.

One of the major problems of the measurement is the accurate determination of these new cavity resonance frequencies and of the line width of the displayed resonance curves between half-power points. In the

solution of this problem, more than ordinary emphasis was placed on ease of operation and elimination of ambiguities in the frequency determination. The resulting instrumentation uses as a stable reference frequency a crystal-controlled oscillator and a frequency multiplier which yields three reference frequencies; 55.8, 223.3 and 4,020 mc. These are mixed in a crystal harmonic generator and mixer leading to a line spectrum of numerous reference frequencies with a constant spacing of 55.8 mc. The same crystal mixer produces a beat frequency signal between the incident signal from the klystron and each of these reference frequencies. A communication receiver with modified IF stage permits selection of a frequency marker out of these beat frequency signals. This marker appears as a blank spot on the traces of incident and reflected signal, and can be moved to any desired point by simply changing the frequency setting of the receiver. Since the receiver dial cannot be read with great accuracy on the high frequency ranges, a frequency counter connected with the local oscillator of the receiver permits a reading of the oscillator frequency to an accuracy of 1 keps. Noting that the local oscillator is 0.455 mc removed from the difference signal, one obtains the frequency of the difference signal with more than sufficient accuracy.

6. Results of Measurements

Transmission- and reflection-type cylindrical cavities have been employed for measurements with linearly and circularly polarized excitation in the 6,000- and 9,000-mc frequency bands. Circular polarization is preferable in the region below saturation where frequency shifts are small.

It is not necessary to choose ferrite discs of relatively small diameter for the purpose of staying within the region of circular polarization close to the cavity axis. The derivation of $\chi_m \pm \kappa$ is not restricted to circularly polarized fields, and the result takes into account that the field becomes more and more elliptically polarized as one approaches the edge of the cavity. This can be seen by rewriting (19) and (20) as follows:

$$\chi_m' \pm \kappa' = \frac{1}{\omega_0 F R_1 R_2} [\Delta\omega_{\pm} (R_1 + R_2) - \Delta\omega_{\mp} (R_1 - R_2)]$$

$$\chi_m'' \pm \kappa'' = \frac{1}{2FR_1R_2} [\Delta(1/Q)_{\pm}(R_1 + R_2) - \Delta(1/Q)_{\mp}(R_1 - R_2)] \quad (26)$$

where $F = t\lambda_0^2/(2L^3)$.

The factor $(R_1 - R_2)$, which is zero for circular polarization, corrects the values for $\chi_m \pm \kappa$ if the disc extends into the region of elliptical

polarization. For relatively small discs ($R_1 = R_2$), one obtains

$$\begin{aligned}\chi_m' \pm \kappa' &= \frac{2}{\omega_0 F R_1} \Delta\omega_{\pm} \\ \chi_m'' \pm \kappa'' &= \frac{1}{F R_1} \Delta(1/Q)_{\pm}\end{aligned}\tag{27}$$

A large number of measurements was made with a half-wave, reflection-type cavity and linearly polarized excitation. Some typical results will be discussed below to demonstrate the applicability of the disc technique. The frequency shift measurements and $\chi_m' \pm \kappa'$ for a ferrite material of 1,300 oersted saturation magnetization are shown on Fig. 6. The ferrite disc has a thickness of 0.0063 inch and completely covers the endwall of the cavity. If the field in the cavity were circularly polarized throughout, then the frequency shift would vary as $\chi_m' \pm \kappa'$. However, elliptical polarization causes a deviation of the measured curves for $\Delta\omega_{\pm}$ from $\chi_m' \pm \kappa'$. Agreement between theoretical and experimental values of $\chi_m' \pm \kappa'$ is good in the regions below and above resonance. Measurements in the resonance region are not possible with a disc of this size because frequency shift and change in Q are so large that the assumption of a small perturbation is violated. In order to establish further that measurements of $\chi_m' \pm \kappa'$ are independent of disc diameter three discs of the same material (saturation magnetization 1300 oersted) with diameters of 0.249, 0.400, and 1.050 inches were measured in the above-mentioned cavity. Plots of χ_m' and κ' (Fig. 7) indicate good agreement for κ' and some scattering of values for χ_m' . This can be explained by noting that the resonance frequency of the empty cavity enters into the computation of χ_m' , but cancels out for κ' (equation 19). Consequently, a very small change in the length of the cavity, as might be expected from opening and reassembling the device, will produce a noticeable error in the low-field region. A change in cavity length of 10^{-4} inch will produce a frequency shift of 1 mc at an operating frequency of 10,000 mc and introduce an error of the order of 0.02 into the measurement of χ_m' . This error may be minimized by using a relatively large disc.

Discs with a diameter of 0.4 inch yielded good measurements of the imaginary quantities $\chi_m'' \pm \kappa''$ in the low-field region. A typical result for a low-loss ferrite (Fig. 8) shows the measured quantities $\Delta(1/Q)_{\pm}$ and the corresponding $\chi_m'' \pm \kappa''$ as a function of the applied magnetic field. (The internal field in the ferrite is obtained by subtracting the magnetization from the applied field). It is noted that only $\Delta(1/Q)_{-}$ can be observed in the resonance region, whereas $\Delta(1/Q)_{+}$ becomes too

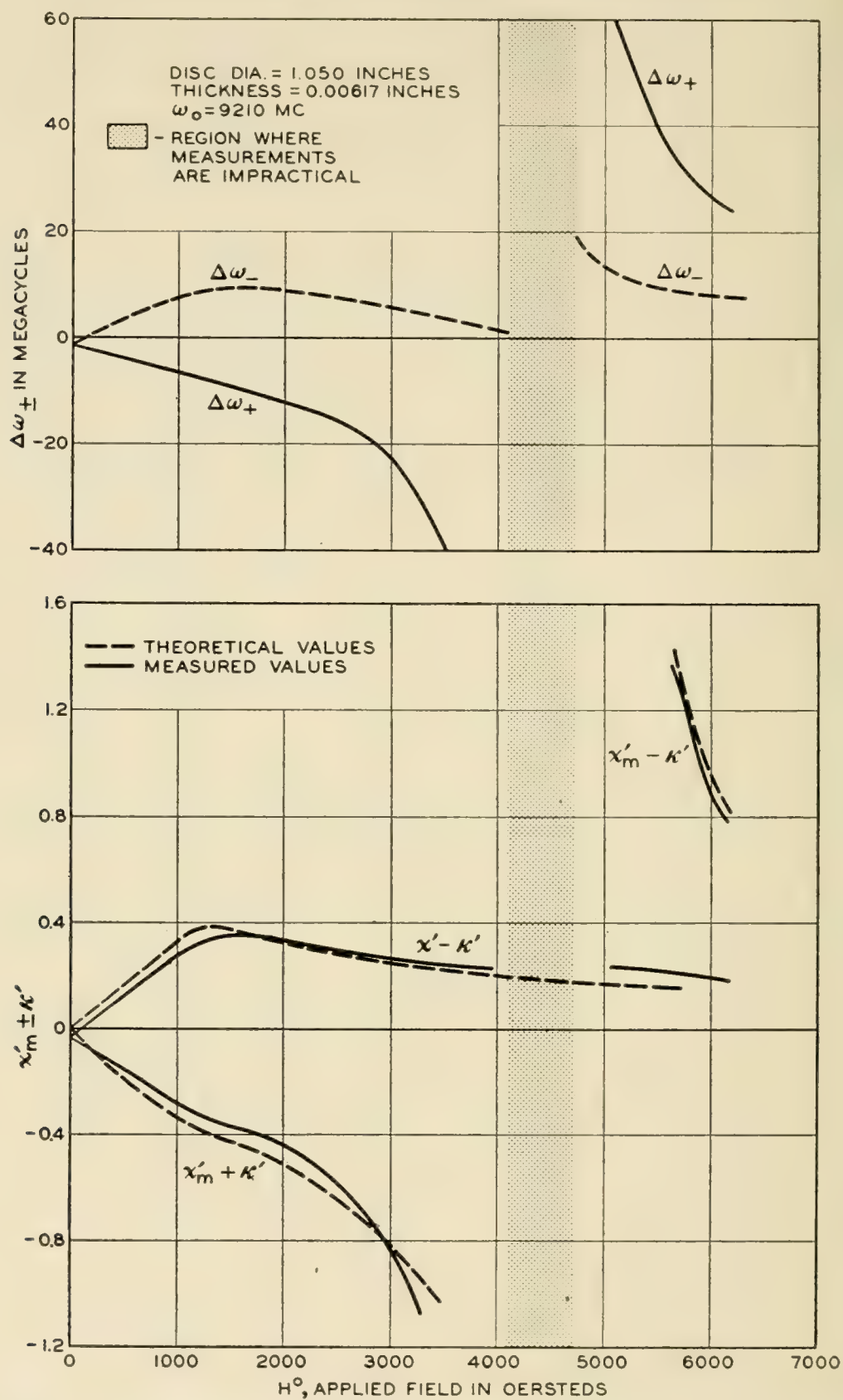


FIG. 6 — Evaluation of $\chi'_m \pm \kappa'$ from measurements of frequency shift $\Delta\omega_{\pm}$ and comparison with Polder's theory. Low-loss BTL ferrite, saturation magnetization 1300 oersted.

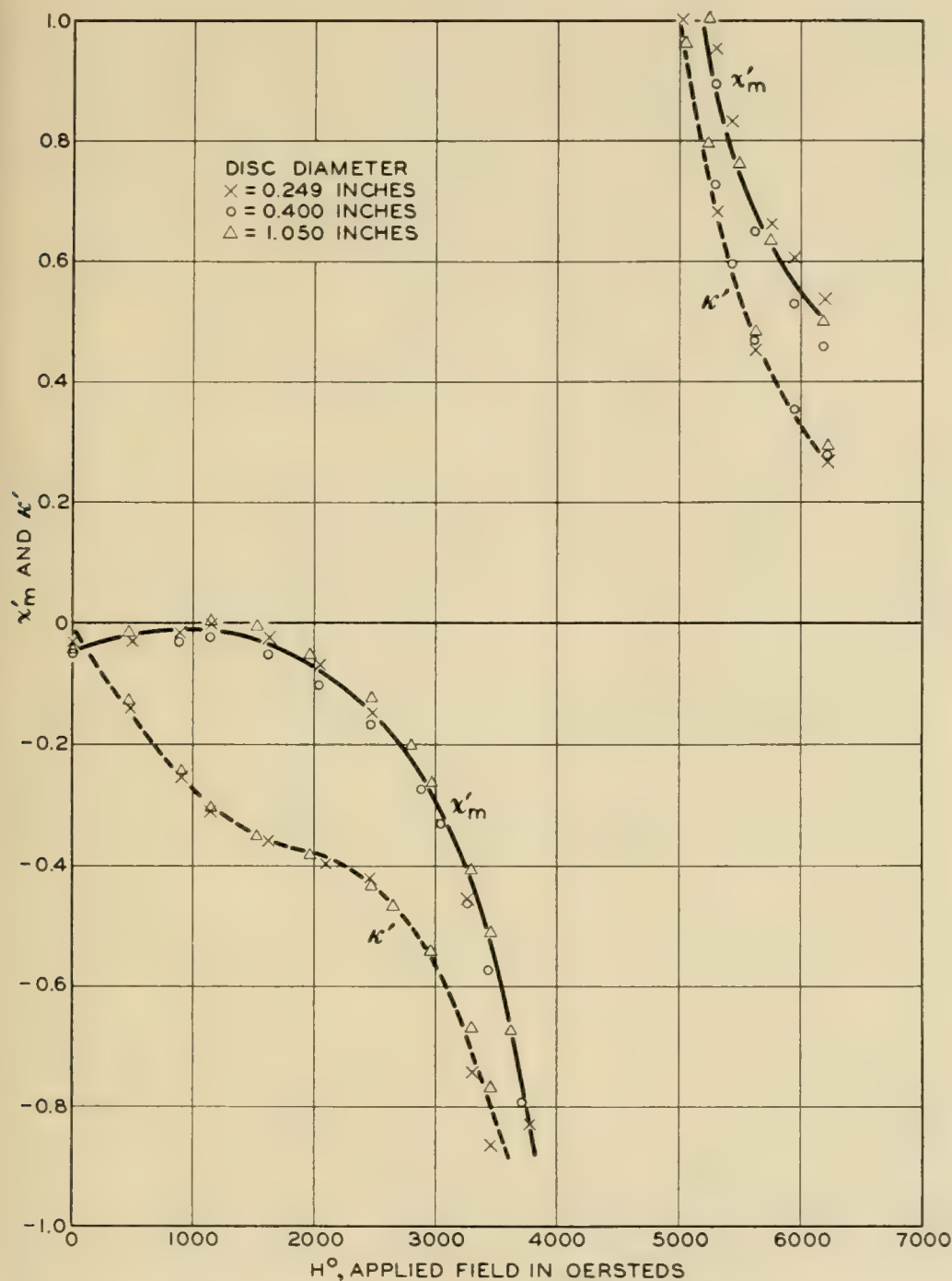


FIG. 7 — Measurement of χ'_m and κ' versus applied static field for three different discs, cut from the same ferrite block.

large to be measured. Knowledge of $\Delta(1/Q)_-$ is not sufficient to determine $\chi''_m - \kappa''$ in the resonance region because here the correction term $\Delta(1/Q)_+(R_1 - R_2)$ (cf. equation 26) becomes comparable to the term $\Delta(1/Q)_-(R_1 + R_2)$ and is needed to compensate for the anomalous peak of the $\Delta(1/Q)_-$ curve. The loss parameters $\chi''_m \pm \kappa''$ assume values of the order of 10^{-2} below and above resonance. The estimated error of the measurement is 3×10^{-3} .

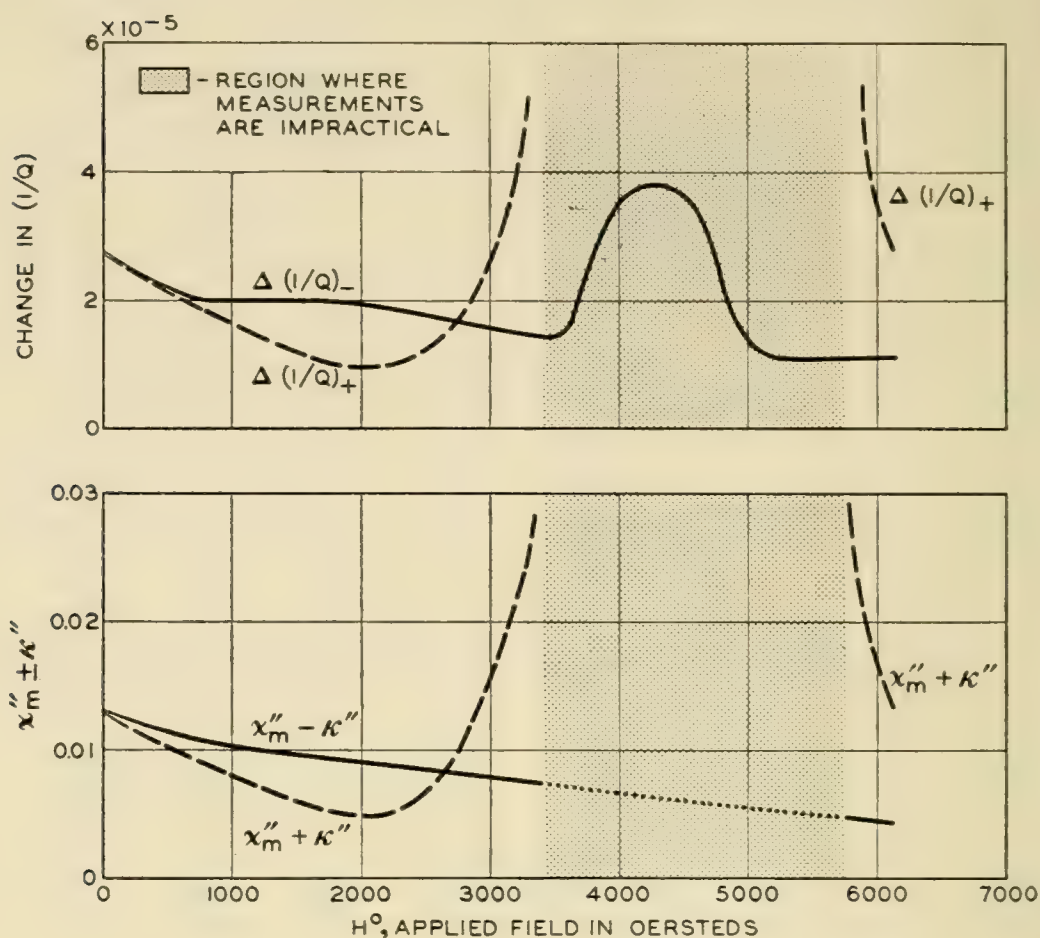


FIG. 8 — Evaluation of $\chi_m'' \pm \kappa''$ from measurements of $\Delta(1/Q)_{\pm}$ for low-loss BTL ferrite, saturation magnetization 1300 oersted.

Whereas it is possible to obtain good agreement between measured and theoretical curves of $\chi_m' \pm \kappa'$ from Polder's relations (3), such an agreement cannot be obtained for $\chi_m'' \pm \kappa''$ from a comparison with (4). This discrepancy is assumed to be caused by the fact that Polder's theory was developed for single ferrite crystals and did not take into account the random orientation of crystal axes in polycrystalline materials, such as the ferrites used in these measurements. It is reasonable to expect a broadening of the resonance line and a departure from the Lorentzian shape in polycrystalline ferrites; hence, the expression for the absorption line (4) no longer holds.

Measurement of the electric susceptibility of ferrites did not present any major difficulties, provided some care was used in the suspension of the discs at the cavity center.

Summing up these results it may be said that the disc method has yielded satisfactory measurements of the intrinsic parameters of polycrystalline ferrites below and above resonance as well as in the un-

saturated region. Measurements of single crystals and of resonance curves with this method have not yet been made.

It is anticipated that in both cases smaller and thinner discs would be needed than have been available so far. However, it can be hoped that these difficulties will be overcome in the near future and that the disc method will be useful also for the study of ferromagnetic resonance phenomena.

APPENDIX

PERTURBATION OF A DEGENERATE CYLINDRICAL CAVITY DUE TO A THIN DISC

The general perturbation equation for a lossless cavity can be derived from energy considerations or directly from Maxwell's equations. We obtain for the shift of resonance frequency due to a small perturbation:

$$2 \frac{\omega_0 - \omega_1}{\omega_0} = \frac{\frac{\mu_0}{2} \int_{v_1} \vec{m} \cdot \vec{h}^{0*} dv + \frac{1}{2} \int_{v_1} \vec{P} \cdot \vec{E}^{0*} dv}{\frac{\mu_0}{2} \int_{v_2} \vec{h}^0 \cdot \vec{h}^{0*} dv} = \frac{W_m^{(s)} + W_e^{(s)}}{W^{(c)}} \quad (\text{A1})$$

ω_0 resonance frequency of the empty cavity

ω_1 resonance frequency of the cavity after insertion of the perturbing sample

$\vec{h}^0$ magnetic field intensity vector in the empty cavity

$\vec{E}^0$ electric field intensity vector in the empty cavity

v_1 volume of sample

v_2 volume of cavity

* indicates the conjugate value

The denominator in (A1) indicates the total energy $W^{(c)}$ stored in the empty cavity at resonance, whereas the numerator is equal to the additional magnetic energy $W_m^{(s)}$ and electric energy $W_e^{(s)}$ stored in the perturbing sample.

Equation (A1) is valid if the frequency shift is small,

$$\frac{\Delta\omega}{\omega_0} = \left| \frac{\omega_1 - \omega_0}{\omega_0} \right| \ll 1 \quad (\text{A2})$$

and if the field in the cavity remains essentially unchanged after insertion of the sample. In order to apply (A1) to the determination of the tensor components χ_m and κ we should attempt to satisfy three conditions: the electric field should vanish at the sample, the magnetic field should be normal to the static magnetic field, and the relationship between RF

magnetization $\vec{m}$ and RF magnetic field $\vec{h}^0$ in the cavity should be simple. All three conditions can be satisfied if a thin ferrite disc is placed against the endwall of the cavity (Fig. 3). Noting that the tangential component of the magnetic field intensity is continuous at the plane face of the disc we have:

$$\begin{aligned} m_r &= \chi_m h_r^0 - j\kappa h_\theta^0 \\ m_\theta &= j\kappa h_r^0 + \chi_m h_\theta^0 \end{aligned} \quad (\text{A3})$$

Inserting (A3) into (A1) we find the additional magnetic energy stored in the disc

$$W_m^{(s)} = \frac{\mu_0}{2} \int_{v_1} [\chi_m (h_r^0 h_r^{0*} + h_\theta^0 h_\theta^{0*}) + j\kappa (h_r^0 h_\theta^{0*} - h_\theta^0 h_r^{0*})] dv \quad (\text{A4})$$

In order to evaluate (A4) we use the fact that the TE₁₁₁-mode can be expressed as the sum of two circularly polarized modes rotating in opposite directions

$$\begin{aligned} \vec{h}^0 &= \vec{h}_1^0 + \vec{h}_2^0 \\ \vec{E}^0 &= \vec{E}_1^0 + \vec{E}_2^0 \end{aligned} \quad (\text{A5})$$

Thus, we have

$$\begin{aligned} h_{r1,2}^0 &= B \frac{\beta}{k_c} J_1'(k_c r) e^{\pm j\theta} \cos \beta z \\ h_{\theta 1,2}^0 &= \pm jB \frac{\beta}{k_c^2 r} J_1(k_c r) e^{\pm j\theta} \cos \beta z \\ h_{z1,2}^0 &= BJ_1(k_c r) e^{\pm j\theta} \sin \beta z \\ E_{r1,2} &= \pm B \frac{\omega\mu_0}{k_c^2 r} J_1(k_c r) e^{\pm j\theta} \sin \beta z \\ E_{\theta 1,2} &= jB \frac{\omega\mu_0}{k_c} J_1'(k_c r) e^{\pm j\theta} \sin \beta z \end{aligned} \quad (\text{A6})$$

(A7)

B	amplitude factor
J_1	Bessel function of the first kind
$\beta = (\beta_0^2 - k_c^2)^{\frac{1}{2}}$	propagation constant in the z -direction
$k_c = p_1'/a$	propagation constant in the r -direction
$\beta_0 = 2\pi/\lambda_0$	propagation constant in free space
L	length of cavity
λ_0	wavelength in free space
a	radius of cavity
$\lambda_g = 2L$	wavelength in the cavity
$p_1' = 1.841$	first zero of the derivative of J_1

Integration of the electric or magnetic field over the volume of the cavity yields the stored energy in the empty cavity at resonance for one of the two circularly polarized modes

$$W^{(c)} = 0.2387 \frac{\pi\mu_0}{4} \cdot \frac{\beta_0^2}{k_c^2} B^2 a^2 L \quad (\text{A8})$$

The magnetic energy $W_m^{(s)}$ in the disc is found by integrating (A4) over the volume of the disc and assuming that the field is constant over the thickness t of the disc. We obtain:

$$W_m^{(s)} = 0.2387 \frac{\pi\mu_0}{2} \frac{\beta^2}{k_c^2} B^2 a^2 t (\chi_m R_1 \pm \kappa R_2) \quad (\text{A9})$$

The two functions R_1 and R_2 depend on the ratio of disc radius to cavity radius

$$R_1 = 4.1893 \frac{r_0^2}{a^2} \left[(J_0(k_c r_0))^2 + \left(1 - \frac{2}{(k_c r_0)^2} \right) (J_1(k_c r_0))^2 \right] \quad (\text{A10})$$

$$R_2 = 2.4720 (J_1(k_c r_0))^2 \quad (\text{A11})$$

It is interesting to note that these two functions are approximately equal (Fig. 4) if the disc radius is less than half the cavity radius. In this region the field in the cavity is essentially circularly polarized, whereas elliptical polarization exists near the wall of the cavity. Inserting (A8) and (A9) into (A1) we find the desired relationship between the two frequency shifts associated with positive and negative circular polarization and the tensor components χ_m and κ .

$$2 \frac{\Delta\omega_{\pm}}{\omega_0} = \frac{1}{2} \frac{\lambda_0^2 t}{L^3} (\chi_m R_1 \pm \kappa R_2) \quad (\text{A12})$$

Equations (A1) and (A12) hold for complex χ_m and κ if a complex frequency shift is introduced as follows:

$$d\bar{\omega} = \omega - \omega_0 + j(\alpha - \alpha_0) \quad (\text{A13})$$

The attenuation constant α may be defined in terms of the internal Q of the cavity, $\alpha = \frac{1}{2}\omega/Q$ and the internal Q is defined as

$$Q = \frac{\omega(\text{energy stored in circuit})}{\text{average power loss}}$$

Thus, the imaginary part of the frequency shift may be expressed as the difference between $(1/Q)$ of the perturbed cavity at the new resonance frequency ω and $(1/Q_0)$ referring to the empty cavity at ω_0 .

$$\Delta(1/Q) = \frac{1}{Q} - \frac{1}{Q_0} = 2 \frac{(\alpha - \alpha_0)}{\omega_0} \quad (\text{A14})$$

We note that the imaginary part of the right hand side of (A1) does indeed represent the power dissipation in the perturbing sample over the stored energy times ω . Hence, taking the imaginary part of (A12)

we have a relationship between the change in $(1/Q)$ and the loss terms χ_m'' and κ''

$$\Delta(1/Q)_{\pm} = \frac{1}{2} \frac{\lambda_0^2 t}{L^3} (\chi_m'' R_1 \pm \kappa'' R_2) \quad (\text{A15})$$

Clearly, (A15) holds only if the initial Q of the empty cavity is very high and the change in Q is quite small.

The electric susceptibility of the ferrite disc can be obtained in a similar way, provided we place the disc at the cavity center where the electric field has a maximum. Then, the additional electric energy stored in the disc is found to be

$$W_e^{(s)} = 0.2387 \frac{\pi}{2} \mu_0 \chi_e \frac{\beta_0^2}{k_c^2} B^2 a^2 t R_1 \quad (\text{A16})$$

It should be noted that there is an important difference between location of a thin disc at the endwalls and at the center of a cylindrical cavity. Whereas the electric field at the endwall may be neglected entirely, the magnetic field at the cavity center has a component parallel to the cavity axis. The effect of this component on the stored energy in the disc may be minimized by magnetizing the disc beyond saturation in the z -direction.

With the assumption that the effects of the magnetic RF field may be neglected we obtain relationships for the electric susceptibility and electric loss factor:

$$2(\Delta\omega/\omega_0) - j\Delta(1/Q) = (\chi_e' - j\chi_e'') \frac{2t}{L} R_1 \quad (\text{A17})$$

Since χ_e is a scalar quantity there is no splitting of the cavity resonance.

ACKNOWLEDGMENT

We would like to thank L. G. Van Uitert who supplied the ferrite materials, Barbara De Hoff who did all of the numerical computation and Edward Kankowski who made most of the measurements shown herein.

REFERENCES

1. D. Polder, *Phil. Mag.*, **40**, p. 99, 1949.
2. C. L. Hogan, *Rev. Mod. Phys.*, **25**, p. 253, 1953.
3. H. Suhl and L. R. Walker, *B.S.T.J.*, **33**, p. 579, 1954.
4. W. A. Yager, J. K. Galt, F. R. Merritt, and E. A. Wood, *Phys. Rev.*, **80**, p. 744, 1950.
5. J. O. Artman and P. E. Tannenwald, *J. Appl. Phys.*, **26**, p. 1124, 1955.
6. A. D. Berk and B. A. Lengyel, *Proc. I.R.E.*, **43**, p. 1587, 1955.
7. A. A. Th. M. Van Trier, *Appl. Sci. Res.*, **3**, p. 305, 1953.
8. J. Stratton, *Electromagnetic Theory*, Chapter 1, McGraw-Hill Book Co., New York, 1941.
9. J. H. Rowen and W. von Aulock, *Phys. Rev.*, **96**, p. 1151, 1954.
10. C. Kittel, *Phys. Rev.*, **73**, p. 155, 1948.

Sensitivity Considerations in Microwave Paramagnetic Resonance Absorption Techniques

By G. FEHER

(Manuscript received February 9, 1956)

This paper discusses some factors which limit the sensitivity of microwave paramagnetic resonance equipments. Several specific systems are analyzed and the results verified by measuring the signal-to-noise ratio with known amounts of a free radical. The two most promising systems, especially at low powers, employ either superheterodyne detection or barretter homodyne detection. A detailed description of a superheterodyne spectrometer is given.

TABLE OF CONTENTS

	Page
I. Introduction	450
II. General Background	450
III. Q Changes Associated with the Absorption	450
IV. Coupling to Resonant Cavities for Maximum Output	451
A. Reflection Cavity	452
1. Detector Output Proportional to Input Power	453
2. Detector Output Proportional to Input Voltage	454
B. Transmission Cavity	455
1. Detector Output Proportional to Input Power	455
2. Detector Output Proportional to Input Voltage	456
V. Minimum Detectable Signal Under Ideal Conditions	457
VI. Signal-to-Noise in Practical Systems	459
A. General Considerations	459
1. Why Field Modulation?	459
2. Choice of Microwave Frequency	460
3. Optimum Amount of Sample to be Used	461
a. Losses Proportional to E^2	461
b. Losses Proportional to H_1^2	462
B. Noise Due to Frequency Instabilities	462
C. Noise Due to Cavity Vibrations	465
D. Klystron Noise	465
E. Signal-to-noise Ratio for Specific Systems	466
1. Barretter Detection	467
a. Straight Detection	469
b. Balanced Mixer Detection	470
2. Crystal Detection	472
a. Simple Straight Detection	473
b. Straight Detection with Optimum Microwave Bucking	473
c. The Superheterodyne Scheme	475
F. Experimental Determination of Sensitivity Limits	477
1. Preparation of Samples	477
2. Comparison of Experimental Result with Theory	478

	Page
VII. A Note on the Effective Bandwidth.....	480
VIII. Saturation Effects.....	482
IX. Acknowledgement.....	483

I. INTRODUCTION

Within the past few years the field of paramagnetic resonance absorption has become an important tool in physical and chemical research. In many ways its usefulness is limited by the sensitivity of the experimental set up. A typical example is the study of semiconductors in which case one would like to investigate as small a number of impurities as possible. It is the purpose of this paper to analyze the sensitivity limits of several experimental set ups under different operating conditions. This was done in the hope that an understanding of these limitations would put one in a better position to design a high sensitivity electron spin resonance equipment. In the last section the performance of the different experimental arrangements is tested. The agreement obtained with the predicted performance proves the essential validity of the analysis. This paper is primarily for experimental physicists confronted with the problem of setting up a high sensitivity spectrometer.

II. GENERAL BACKGROUND

We will not consider here the detailed theory¹ of the resonance phenomenon but consider this part of the problem only from a phenomenological point of view. When a paramagnetic sample is placed into an RF field of amplitude H_1 of a frequency ω at right angles to which there is a dc magnetic field H_0 , magnetic dipole transitions will be induced in the neighborhood of the resonance condition

$$h\omega = g\beta H_0 \quad (1)$$

where g is the spectroscopic splitting factor, h is Planck's constant and β is the Bohr magneton. As a result of these transitions power will be absorbed from the microwave field H_1 . This power absorption is associated with the imaginary part of the RF susceptibility χ'' . The transmitted (or reflected) H_1 will also experience a phase shift which is associated with the real part of the RF susceptibility χ' . The sensitivity of the setup is then determined by how small a power absorption (or phase shift) one is able to detect when going through a resonance

III. Q CHANGES ASSOCIATED WITH THE ABSORPTION

The average power absorbed per unit volume of a paramagnetic sample is

$$P = \frac{1}{2}\omega H_1^2 \chi'' \quad (2)$$

For low enough powers χ'' is not a function of H_1 (and even for very high powers never drops off faster than $1/H_1^2$), so that for a large power absorption one would like a large RF magnetic field. This suggests a resonant cavity which indeed is used in all experimental setups. The Q of a cavity into which a paramagnetic sample is placed is given by

$$Q = \omega \frac{\text{Energy Stored}}{\text{Average Power Dissipated}} = \omega \frac{\frac{1}{8\pi} \int_{V_c} H_1^2 dV_c}{P_1 + \frac{1}{2} \omega \int_{V_s} H_1^2 \chi'' dV_s} \quad (3)$$

where P_1 = power dissipated in the cavity in the absence of any paramagnetic losses, V_s is the sample volume and V_c the cavity volume.

Assuming that the paramagnetic losses are small in comparison with P_1 we get

$$Q = Q_0 \left(1 - 4\pi \frac{\int_{V_s} H_1^2 \chi'' dV_s}{\int_{V_c} H_1^2 dV_c} Q_0 \right) = Q_0 (1 - 4\pi \chi'' \eta Q_0) \quad (4)$$

$$\therefore \Delta Q = Q_0^2 4\pi \chi'' \eta$$

where Q_0 is the cavity Q in the absence of paramagnetic losses and η is the filling factor and depends on the field distribution in the cavity and the sample. For example, in a rectangular cavity excited in the TE_{101} mode

$$\eta = \frac{V_s}{V_c} \frac{4}{1 + \left(\frac{d}{a}\right)^2} \quad (5)$$

where d is the length of the cavity and a the width along which the E field varies. In the above example it was assumed that the sample is small in comparison to a wavelength and is placed in the max. H_1 field.

IV. COUPLING TO RESONANT CAVITIES FOR MAXIMUM OUTPUT

Having established the Q changes associated with the resonance absorption, we will next determine the proper coupling to the resonant cavity in order that the Q changes result in a maximum change in transmitted or reflected power (or voltage). The derivation will be based on the assumption that we have a fixed amount of power available from our source and that the Q change is not a function of the RF power (no saturation effects).

A. Reflection Cavity

Fig. 1 shows a magic (hybrid) T which serves to observe the reflected power from the cavity. Arm 3 has a slide screw tuner which serves to balance out some of the power coming from arm 2. This does not affect the present analysis and will be considered later in connection with detector noise. It should be mentioned, however, that a certain amplitude or phase unbalance has to be left. This insures that the signal in arm 4 will be a function of either χ' or χ'' .¹ In the case that the magic T is completely balanced out the signal in arm 4 will be a function of both χ' and χ'' and the experimental results become difficult to analyze.

Fig. 2 shows the equivalent circuit for a reflection cavity.² The $\sqrt{2}$ in the source voltage arises from the fact that half the power is lost in arm 3. From this equivalent circuit we can define the following relations:

$$\text{Unloaded } Q = Q_0 = \frac{\omega L}{r} \quad (\text{Losses due to cavity alone}) \quad (6)$$

$$\text{External } Q = Q_x = \frac{\omega L}{R_0 n^2} \quad (\text{Losses arising from power } R_0 n^2 \text{ leaking out of the cavity}) \quad (7)$$

$$\text{Loaded } Q = Q_L = \frac{\omega L}{R_0 n^2 + r} \quad (\text{Losses due to both cavity } R_0 n^2 \text{ and leakage out}) \quad (8)$$

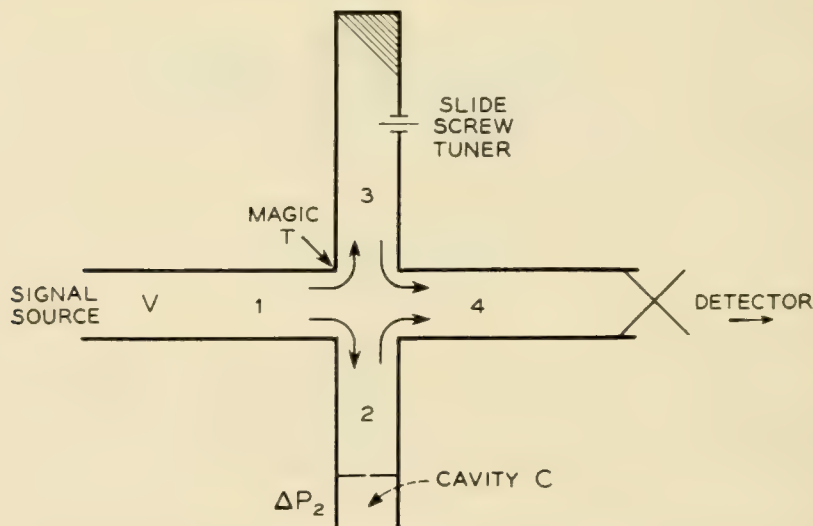


Fig. 1 — A simple arrangement to observe the reflected power from cavity C.

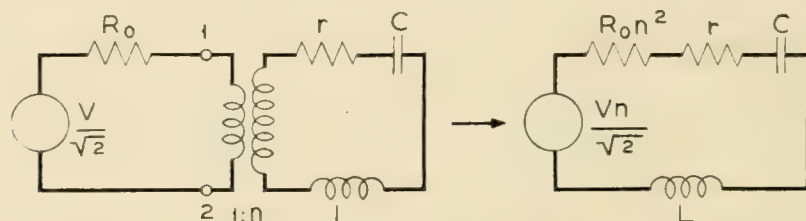


Fig. 2 — Equivalent circuit for a reflection cavity.

We define the coupling coefficient $\beta = Q_x/Q_0$ such that:

$$\text{Critically coupled cavity } \beta = \frac{Q_0}{Q_x} = \frac{R_0 n^2}{r} = 1 \quad (9)$$

$$\text{Overcoupled cavity } \frac{Q_0}{Q_x} > 1 \quad \text{VSWR} = \beta = \frac{R_0 n^2}{r} \quad (10)$$

$$\text{Undercoupled cavity } \frac{Q_0}{Q_x} < 1 \quad \text{VSWR} = \frac{1}{\beta} = \frac{r}{R_0 n^2} \quad (11)$$

We will see that the coupling coefficient for maximum output depends on the characteristics of the detecting element. Two cases will be treated: the power (square law) detector and the voltage (linear) detector.

1. Detector Output Proportional to Incident Power

Power into the cavity at resonance

$$P_c = \left(\frac{Vn}{\sqrt{2}} \right)^2 \frac{r}{(R_0 n^2 + r)^2}$$

Max. power available from source (in arm 1)

$$P_0 = \frac{(Vn)^2}{4R_0 n^2} \therefore P_c = P_0 \frac{2R_0 n^2 r}{(R_0 n^2 + r)^2}$$

The change in reflected power ΔP_R equals the change in the power inside the cavity ΔP_c (since the incident power stays the same).

$$\Delta P_c = \frac{\partial P_c}{\partial r} \Delta r = 2R_0 n^2 P_0 \frac{R_0 n^2 - r}{(R_0 n^2 + r)^3} \Delta r \quad (12)$$

We want to optimize ΔP_c with respect to the coupling parameter n^2 (or $R_0 n^2$), i.e.,

$$\begin{aligned} \frac{\partial(\Delta P_c)}{\partial(R_0 n^2)} &= (R_0 n^2)^2 - 4(R_0 n^2)r + r^2 = 0 \\ \therefore \frac{R_0 n^2}{r} &= 2 \pm \sqrt{3} \end{aligned} \quad (13)$$

the positive sign being associated with the overcoupled, the negative with the undercoupled case. The experimentally measured quantity is the voltage standing wave ratio $\text{VSWR} = 2 + \sqrt{3} = 3.74$ corresponding to a reflection coefficient of 0.58. Putting this value into (12) we get for the maximum signal

$$\frac{\Delta P_c}{P_0} = \pm 0.193 \frac{\Delta r}{r} = \mp 0.193 \frac{\Delta Q_0}{Q_0} = \mp (0.193)(4\pi)\chi''\eta Q_0 \quad (14)$$

the last step being obtained with the aid of (4). Equation (12) is plotted in Fig. 4. From the symmetry of the graph it is obvious that for a given VSWR, the signal will be the same for the overcoupled and undercoupled case. However, as we will see later from the standpoint of noise the 2 cases are not necessarily identical.

2. Detector Output Proportional to Input Voltage

Let Γ be the reflection coefficient, then V_{REFL} from the cavity is

$$V_{\text{REFL}} = \frac{V}{\sqrt{2}} \Gamma = \frac{V}{\sqrt{2}} \left(\frac{\text{VSWR} - 1}{\text{VSWR} + 1} \right) = - \frac{V}{\sqrt{2}} \left[1 - \frac{2\text{VSWR}}{\text{VSWR} + 1} \right]$$

With the aid of (10) and (11) this gives for the undercoupled case:

$$V_{\text{REFL}} = - \frac{V}{\sqrt{2}} \left(1 - \frac{2r}{R_0 n^2 + r} \right)$$

and the overcoupled case

$$V_{\text{REFL}} = - \frac{V}{\sqrt{2}} \left(1 - \frac{2R_0 n^2}{R_0 n^2 + r} \right)$$

We are interested only in the change of output voltage which is:

$$\Delta V_{\text{REFL}} = \frac{\partial V_{\text{REFL}}}{\partial r} \Delta r = \pm \sqrt{2} V \Delta r \frac{R_0 n^2}{(R_0 n^2 + r)^2} \quad (15)$$

The two signs corresponding to the undercoupled or overcoupled case, respectively. In order to find the optimum coupling

$$\begin{aligned} \frac{\partial(\Delta V)}{\partial(R_0 n^2)} &= R_0 n^2 - r = 0 \\ \therefore \frac{R_0 n^2}{r} &= 1 \end{aligned}$$

Putting this value into (15) we get the max. value

$$\frac{\Delta V_{\text{REFL}}}{V} = \pm \frac{\sqrt{2}}{4} \frac{\Delta r}{r} = \mp \frac{\sqrt{2}}{4} \frac{\Delta Q_0}{Q_0} = \mp \frac{\sqrt{2}}{4} 4\pi \chi'' \eta Q_0 \quad (16)$$

(15) is again plotted in Fig. 4. From this graph we see that for maximum sensitivity we want to work near match. However, one should not work so close to match that the absorption signal will carry the cavity through the matching condition while sweeping through a resonance line. This would result (due to the sign reversal of the signal at match) in a distorted line. Incidentally, the sign of the signal may be conveniently used

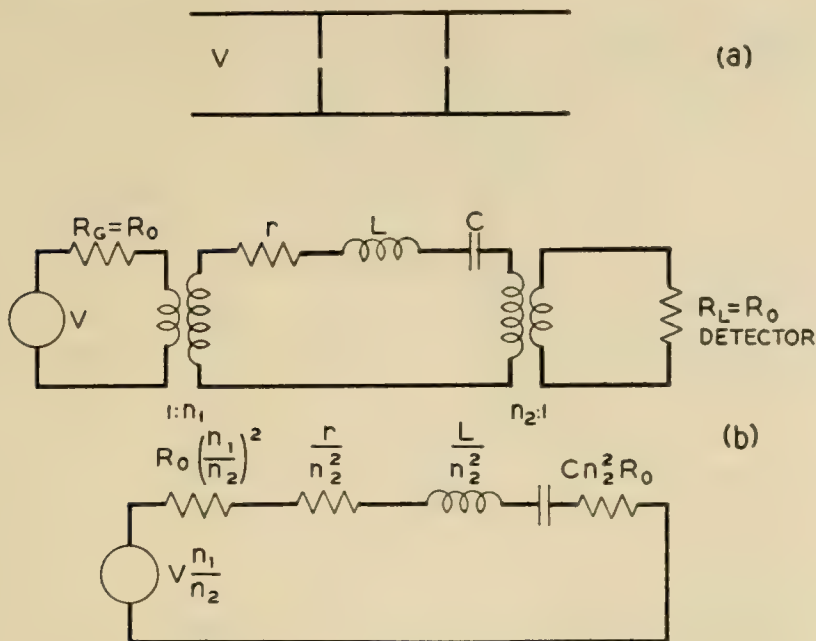


Fig. 3 — Equivalent circuit for a transmission cavity.

to determine whether the cavity is overcoupled or undercoupled. This information may be necessary in Q_0 determinations.¹²

B. Transmission Cavity

Fig. 3 shows the equivalent circuit for a transmission cavity,² the generator and the detector being matched to the waveguide, i.e.,

$$R_G = R_L = R_0$$

Analogous to the reflection cavity we define again two coupling coefficients

$$\beta_1 = \frac{R_0 n_1^2}{r} \quad \beta_2 = \frac{R_0 n_2^2}{r} \quad (18)$$

the relation between the unloaded and loaded Q being

$$Q_0 = Q_L(1 + \beta_1 + \beta_2) \quad (19)$$

1. Detector Output Proportional to Input Power

$$\text{Power into load } P_L = \frac{V^2 n_1^2 n_2^2 R_0}{(R_0 n_1^2 + r + R_0 n_2^2)^2}$$

$$\text{Max. power generator can deliver } P_0 = \frac{(V n_1)^2}{4(n_1^2 R_0)} \quad (19)$$

$$\Delta P_L = \frac{\partial P_L}{\partial r} \Delta r = \frac{2V^2 n_1^2 n_2^2 R_0}{(R_0 n_1^2 + r + R_0 n_2^2)^3} \Delta r$$

In the case of the transmission cavity we have two coupling coefficients whose optimum value we have to determine.

$$\frac{\partial(\Delta P_L)}{\partial(n_1^2 R_0)} = \frac{\partial(\Delta P_L)}{\partial(n_2^2 R_0)} = 0 \quad (20)$$

$$\therefore n_1^2 R_0 = n_2^2 R_0 = r$$

which means that the input and output coupling should be identical. Relation 20 looks superficially like a matching condition. However, it should be noted that the input impedance to the cavity contains besides the cavity impedance the load impedance. Hence, the VSWR is

$$\frac{R_0 n^2 + r}{R_0 n^2}$$

which represents an undercoupled case. One never can overcouple a transmission cavity with equal input and output couplings. Putting condition (20) into (19) we get:

$$\frac{\Delta P_L}{P_0} = -\frac{8}{27} \frac{\Delta r}{r} = -\frac{8}{27} \frac{\Delta Q_0}{Q_0} = -\left(\frac{8}{27}\right) 4\pi\chi''\eta Q_0 \quad (21)$$

2. Detector Output Proportional to Input Voltage

The voltage across the load

$$V_L = \frac{V n_1 n_2 R_0}{R_0 n_1^2 + r + R_0 n_2^2}$$

Again for max. sensitivity both couplings should be the same

$$\Delta V_L = \frac{\partial V_L}{\partial r} \Delta r = -V \left[\frac{n_2^2 R_0}{(r + 2R_0 n_2^2)^2} \right] \Delta r \quad (22)$$

$$\frac{\partial V_L}{\partial(n_2^2 R_0)} = 0 \quad \therefore R_0 n^2 = \frac{r}{2} \quad (23)$$

$$\therefore \frac{\Delta V_L}{V} = -\frac{1}{8} \frac{\Delta r}{r} = -\frac{1}{8} \frac{\Delta Q_0}{Q_0} = -\frac{1}{8} 4\pi\chi''\eta Q_0 \quad (24)$$

Fig. 4 is a plot of (12), (15), (19), and (22). It should be noted that the sensitivities of the reflection cavity are normalized to the input of the magic T (in Fig. 1) and not to the input of the cavity as in the transmission cases. This results in a 3-db decrease in output and causes the power sensitivity of the transmission cavity to look relatively higher. However this is somewhat arbitrary since a balanced transmission type scheme would also require a magic T with an accompanying reduction in usable power.

All the previous sensitivity expressions are proportional to χ'' . For an unsaturated condition it may be replaced by χ' whenever the output is sensitive to phase changes in the cavity.¹

It should be noted that we maximized the output from the detector. This will result in a maximum signal to noise ratio if the noise is a constant independent of the microwave power. This, however, is in general not the case and in the next section we will investigate the signal to noise ratio taking into account its dependence on the RF power.

V. MINIMUM DETECTABLE SIGNAL UNDER IDEAL CONDITIONS

The minimum signal is ultimately determined by the random thermal agitation. Due to this cause the power fluctuates by an amount $kT\Delta\nu$, where k is Boltzmann's constant, T the absolute temperature and $\Delta\nu$ the bandwidth. The minimum detectable microwave power will be then of the order of $kT\Delta\nu$. This is the problem one faces when designing sensitive microwave receivers. However, our problem is of a different nature. We want to detect a small change in the power level of a relatively large

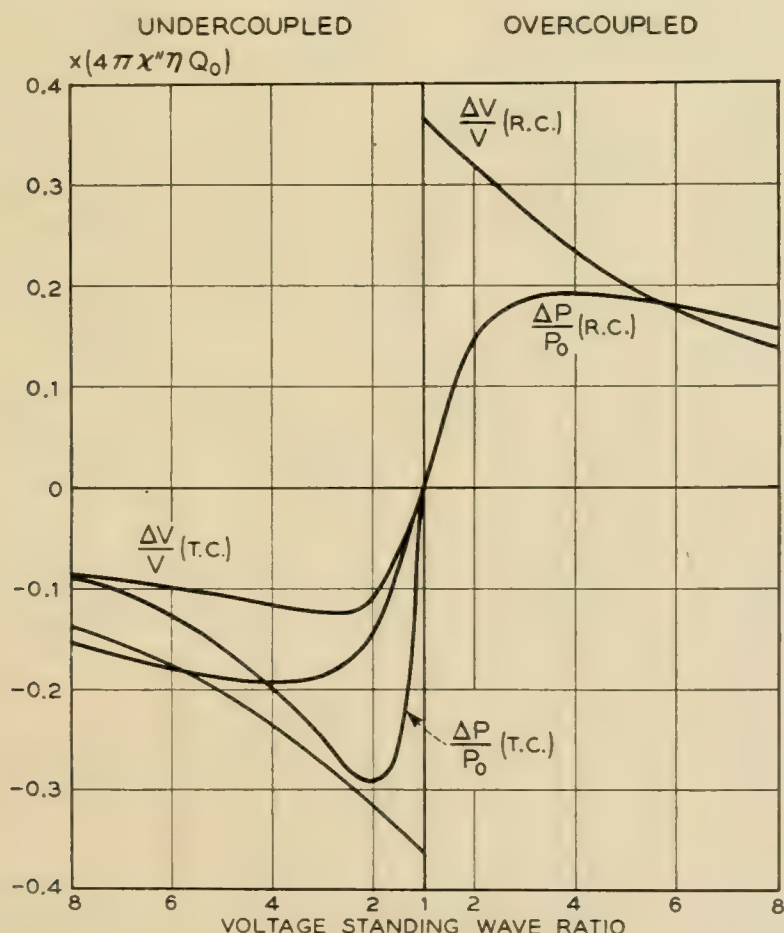


Fig. 4 — Output versus V.S.W.R. for reflection and transmission cavity.

microwave signal. This change in power level will have to be considerably larger than $kT\Delta\nu$ before it can be detected.³ The physical reason for this is that the fluctuating fields associated with the noise power combine with the microwave fields to produce power fluctuations much larger than $kT\Delta\nu$. It is more straightforward to compare noise voltages rather than powers, especially since the power changes are not necessarily a constant of the system. In Fig. 1 for instance, the power change in arm 4 is

$$\Delta P_4 = \frac{(\Delta V)(V_4)}{R_0}$$

whereas the power change in arm 2 (for the same voltage change) is

$$\Delta P_2 = \frac{(\Delta V)(V_2)}{R_0}$$

It also shows that one wants to maximize the change in output voltage as was done in Section IV.

The open terminal RMS noise voltage of a system with an internal impedance R_0 is given by

$$V_{\text{RMS}} = \sqrt{4R_0kT\Delta\nu}$$

If we terminate this system with a noiseless resistor R_0 , the voltage across it will be $\sqrt{R_0kT\Delta\nu}$. However, the terminating resistor is also at temperature T , so that the total RMS voltage across it will be $\sqrt{2} \sqrt{R_0kT\Delta\nu}$.

Comparing this RMS noise voltage with the signal voltage obtained in (16), we get for the reflection cavity*

$$\Delta V = V \sqrt{2\pi\chi''} \eta Q_0 = \sqrt{2} \sqrt{R_0kT\Delta\nu} \quad (25)$$

$$\therefore \chi_{\text{MIN}}'' = \frac{1}{Q_0\eta\pi} \left(\frac{kT\Delta\nu}{2P_0} \right)^{\frac{1}{2}} \quad (26)$$

As an example let us consider the following typical value for a 3-cm setup. $Q_0 = 5 \times 10^3$, $\Delta\nu = 0.1$ cps, $P_0 = 10^{-2}$ Watts

$$\eta = \frac{V_s}{V_c} \frac{4}{1 + \left(\frac{d}{a}\right)^2} \simeq \frac{4V_s}{10 \text{ cm}^3}.$$

For this case

$$(\chi_{\text{min}}'')(V_s) = \simeq 2 \times 10^{-14}$$

* In most cases the behaviour of the transmission and reflection cavity is similar, so that they will not be treated separately.

This corresponds for an unsaturated Lorentz line¹ to a static susceptibility $\chi_0 = \sqrt{3}\chi''(\Delta\omega/\omega)$, where $\Delta\omega$ is the line width between inflection points. For the free radical diphenyl picryl hydrazyl having a 2 oersted line width this expression at room temperature gives for the min. number of spins 10^{10} . A plot of the minimum RF susceptibility and minimum number of electrons versus microwave power is shown in Fig. 5.

VI. SIGNAL-TO-NOISE IN PRACTICAL SYSTEMS

A. General Considerations

1. Why Field Modulation?

From a design point of view it is instructive to consider the minimum fractional voltage change corresponding to the above $\chi_{\min}''V_s$ of 2×10^{-14} . This turns out to be, see (16),

$$\frac{\Delta V_{\min}}{V} \simeq 2 \times 10^{-10}.$$

From this figure one may safely conclude that it is not feasible to use any system in which the microwave carrier level reflected from the cavity has to be kept constant to this accuracy. Such systems would include straight detection, the dc being bucked out and amplified or systems employing amplitude modulation of the carrier. (Although

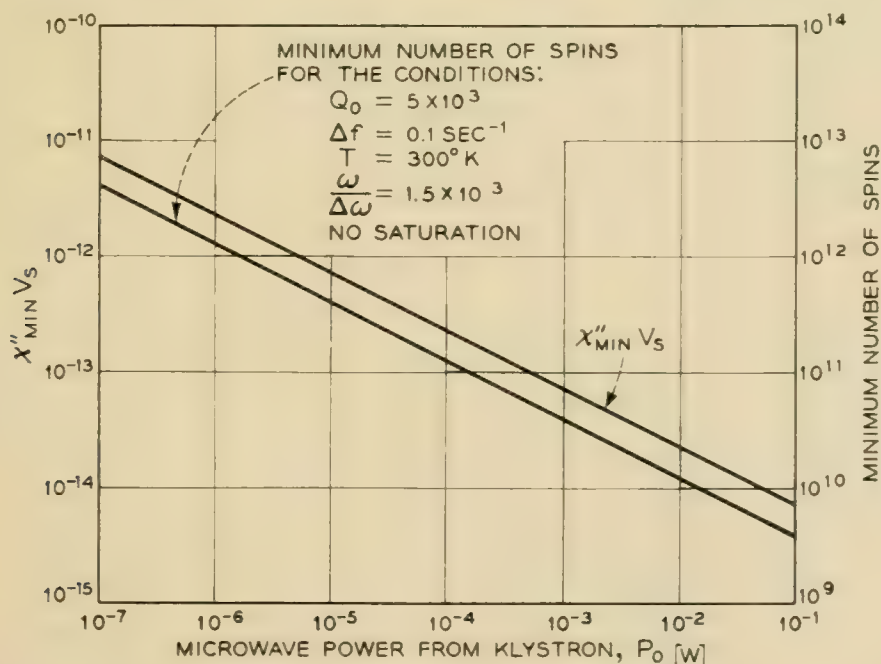


Fig. 5 — Minimum RF susceptibility and number of electrons which should be observable under thermal noise limitations. The conditions for the minimum number of spins correspond closely to those under which the experimental set-ups were tested.

the latter system may be improved by microwave bucking, it still remains very much inferior to the field modulation system to be described presently.) A system commonly used in which the requirements on the constancy of the microwave level is less stringent, makes use of a small external magnetic field modulation of angular frequency ω_M superimposed on the slowly varying dc magnetic field. Thus in the absence of a resonance line the output is zero except for some small Fourier components of the random fluctuations at the frequency ω_M . If the amplitude of the field modulation ΔH_M is small in comparison to the line width ΔH this method will sweep out the derivative of the line, i.e., the signal will not be proportional to χ'' as previously assumed but to $d\chi''/dH (\Delta H_M)$. In order to preserve the line shape one should sweep only over a fraction of the line width. The sensitivity will thereby be reduced by roughly the same fraction. It should be noted, however, that even if one overmodulates the line (in order to increase the sensitivity) the resonance condition (i.e., place of zero signal, corresponding to new slope in the absorption) will not shift for a symmetrical line and the correct g -value may be obtained. Also from the knowledge of the amplitude of the modulating field the increase in line width may be corrected for. For those reasons we will not be concerned with the reduction in sensitivity due to this field modulation scheme.

2. Choice of Frequency

Referring to (26)

$$\chi_{\min} \propto \left(\frac{V_c}{V_s} \right) \frac{1}{Q_0} \frac{1}{\sqrt{P_0}}$$

and the minimum total number of electrons $N_{\min}$.

$$N_{\min} \propto \chi_0 V_s \propto \left(\frac{V_c}{Q_0} \right) \left(\frac{\Delta\omega}{\omega} \right) \frac{1}{\sqrt{P_0}} \quad (27)$$

Assuming that we are dealing with the same type of cavity mode at different frequencies, the same power, and that the line width $\Delta\omega$ is constant, we have $V_c \propto \frac{1}{\omega^3}$, $Q_0 \propto \frac{1}{\omega^{\frac{1}{2}}}$

$$\therefore N_{\min} \propto \frac{1}{\omega^{7/2}} \quad (28)$$

Equation 28 shows that in order to see the smallest number of spins we want to go to as high as frequency as possible. The upper limit is given by the availability of components in the millimeter region, by the difficulty of handling them and by the maximum available power. The most commonly used setups operate at a wavelength of 1 cm and 3 cm.

The latter was used in the experimental part of this paper. From (28) we see that with a 1 cm setup the number of observable electrons should be approximately 40 times less than with a 3-cm setup. However, in most practical cases one is not limited by the amount of available sample, since one usually can increase the sample size at longer wavelengths. Therefore a better criterion is the minimum number of electrons per unit volume

$$\frac{N_{\min}}{V_s} \propto \left(\frac{V_c}{V_s}\right) \frac{1}{Q_0} \left(\frac{\Delta\omega}{\omega}\right) \quad (29)$$

Keeping now the filling factor V_c/V_s constant we see that the sensitivity of a 3 cm setup as defined in (29) is only $\sqrt{3}$ worse than of a 1 cm setup. In addition the power outputs of 3-cm klystrons are usually sufficiently higher than those of 1-cm klystron to overcome even the $\sqrt{3}$ advantage. If RF saturation comes in, the power argument is not valid, but one has to consider the RF magnetic field H_1 inside the cavity which for a given power and Q_0 is prop. to ω . Thus at the higher frequencies (1 cm) saturation effects become more pronounced reducing again the advantage of a 1-cm over a 3-cm setup. In deciding the choice of the frequency in special cases (e.g., when $\Delta\omega$ is a function of the magnetic field, or the sample is larger than a skin depth) (29) should be used.

There are, of course, considerations, other than those of max. sensitivity, which have to be taken into account. For example one would always like to satisfy the condition $\Delta\omega/\omega \ll 1$ which favors higher frequencies. On the other hand, for very narrow lines (say less than 0.1 oersteds) fractional field instabilities and inhomogeneities will favor low magnetic fields, i.e., lower frequencies. One also might encounter samples which exhibit an excessive loss in a given frequency band which therefore has to be avoided.

3. Optimum Amount of Sample to be Used

The output voltage ΔV is proportional to the sample volume and the unloaded Q_0 of the cavity, see (26). If the sample is lossy an increase in its size will reduce the Q and therefore reduce the signal. We may roughly distinguish two limiting cases. In one case the losses are proportional to E^2 (e.g., high resistivity samples having dielectric losses), in the other case they are proportional to H_1^2 (e.g., low resistivity samples in which the losses are due to surface currents).

a. Losses Proportional to E^2

The paramagnetic sample is placed in the region of max. RF magnetic field for instance at the end plate of a rectangular cavity resonating in

the TE_{10} mode. The sample extends a distance x into the cavity, x being assumed to be small in comparison to the wavelength so that H_1 does not vary appreciably over the sample. The additional losses due to the sample will then be proportional to

$$A \int_0^x E_0^2 \sin^2 \left(\frac{2\pi x}{\lambda} \right) dx \sim Cx^3$$

whereas the volume of the sample is proportional to x . The observed voltage change ΔV is then given by

$$\Delta V \propto Q_0' V_s \propto \left(\frac{1}{\frac{1}{Q_0} + Cx^3} \right) (x)$$

where Q_0 is the Q of the cavity without sample, and Q_0' the total Q of both cavity and sample. Maximizing the voltage change ΔV we get

$$\begin{aligned} \frac{\partial(\Delta V)}{\partial x} = 0 \quad \therefore Cx^3 &= \frac{1}{2Q_0} \\ \therefore Q_0' &= \frac{2}{3}Q_0 \end{aligned} \quad (30)$$

Equation (30) tells us that we should load the cavity with the sample until the Q_0 is reduced to $\frac{2}{3}$ of its original value. It should be pointed out that we optimized the signal and not the more important quantity, the signal-to-noise ratio. Therefore the analysis is only valid as long as the noise is not a function of Q_0 , (see Section VIB) and that we do not saturate the sample (see Section VIII). If either condition does not hold the amount of sample to be put in should exceed the above calculated value.

b. Losses Proportional to H_1^2

The losses do not vary along the sample so that we may write

$$\Delta V \propto Q_0' V_s \propto \left(\frac{1}{\frac{1}{Q_0} + Cx} \right) (x)$$

which clearly has no maximum for x . One should therefore put as big a sample into the cavity as possible (compatible with the assumption that if it be small in comparison to a wavelength), the same result as if one had no losses at all.

B. Noise Due to Frequency Instabilities

Before considering the signal-to-noise ratio for specific systems we will investigate a noise source which is common to all of them. It arises

from the random frequency variations of the microwave source or from the random variations of the resonant frequency of the cavity (e.g., rising helium bubbles at 4°K or just any microphonics).

From the equivalent circuit of a reflection cavity (see Fig. 2) we can write for the voltage standing wave ratio

$$\begin{aligned} \text{VSWR} = \frac{Z_{in}}{R_0 n^2} &= \frac{r}{R_0 n^2} + \frac{j}{R_0 n^2} \left(\omega L - \frac{1}{\omega C} \right) \\ &= \frac{r}{R_0 n^2} \left[1 + jQ_0 \left(\frac{2\Delta\omega}{\omega} \right) \right] = \frac{r}{R_0 n^2} [1 + j\delta] \end{aligned}$$

where $\delta = Q_0(2\Delta\omega/\omega)$ and $\Delta\omega$ is the frequency deviation from the resonant frequency of the cavity. The reflection coefficient Γ is then given by:

$$\Gamma = \frac{\frac{r}{R_0 n^2} (1 + j\delta) - 1}{\frac{r}{R_0 n^2} (1 + j\delta) + 1} = \Gamma_0 + 2 \frac{\frac{R_0 n^2}{r} \delta^2}{\left(\frac{R_0 n^2}{r} + 1 \right)^3} + 2j \frac{\frac{R_0 n^2}{r} \delta}{\left(\frac{R_0 n^2}{r} + 1 \right)^2}$$

where Γ_0 is the reflection coefficient at the resonant frequency of the cavity. The other two terms giving the changes in Γ for a given frequency deviation $\Delta\omega$. The changes of $\Delta\omega$ (or $\Delta\Gamma$) having ac components near the modulation frequency will thus represent noise terms which will pass together with the signal through the detection system. The slide screw tuner (see Fig. 1) which is used to buck out part of the microwave power will introduce an additional reflection coefficient $\Gamma_R + j\Gamma_R'$. Thus the total reflection coefficient will be given by

$$\Gamma = \Gamma_0 - \Gamma_R + 2 \frac{\frac{R_0 n^2}{r} \delta^2}{\left(\frac{R_0 n^2}{r} + 1 \right)^3} - j\Gamma_R' + j2 \frac{\frac{R_0 n^2}{r} \delta}{\left(\frac{R_0 n^2}{r} + 1 \right)^2} \quad (31)$$

We are interested only in the magnitude of V (i.e., $|\Gamma|$) reaching the detector. Tuning to the dispersion mode (χ') the slide screw tuner is adjusted such that $\Gamma_R' \gg \Gamma_0 - \Gamma_R$. Under these conditions the output noise voltage will be given by

$$\left(\frac{\Delta V_N}{V} \right)_{\chi'} \simeq 2 \frac{\frac{R_0 n^2}{r} \delta}{\left(\frac{R_0 n^2}{r} + 1 \right)^2} \quad (32)$$

Tuning to the absorption (χ''), the condition $\Gamma_0 - \Gamma_R \gg \Gamma_R$ will be

satisfied and (31) becomes

$$\left(\frac{\Delta V_N}{V}\right)_{\chi''} \simeq 2 \frac{\left(\frac{R_0 n^2}{r}\right) \delta^2}{\left(\frac{R_0 n^2}{r} + 1\right)^3} + 2 \frac{\left(\frac{R_0 n^2}{r}\right)^2 \delta^2}{\left(\frac{R_0 n^2}{r} + 1\right)^4 (\Gamma_0 - \Gamma_R)} \quad (33)$$

An inspection of (31), (32), and (33) shows that the noise voltage, enters as a first order effect in δ when tuned to χ' . This is not too surprising since a frequency effect is expected to affect predominantly the dispersion mode. When tuned to χ'' the effect becomes second order as long as $|\Gamma_0 - \Gamma_R|$ is large. Under those conditions the 2 terms in (33) are of comparable magnitude. We can easily see the origin of the second term. It arises from the first order out-of-phase component of the noise voltage. Being, however, sensitive only to in-phase components it will be reduced to a second order effect—as long as $|\Gamma_0 - \Gamma_R|$ is large, i.e., as long as we have a carrier which makes us insensitive to out-of-phase components.* When $\Gamma_0 - \Gamma_R$ goes to zero (33) ceases to hold and the noise voltage will be given by (32).

There are two important conclusions to be drawn from (33).

We want to keep $\Gamma_0 - \Gamma_R$ as large as possible. Therefore in schemes (like the superheterodyne see section VI E) where this is not feasible, special care has to be taken to eliminate this noise source.

From (15) we find that the desired signal is proportional to

$$\frac{R_0 n^2}{r} \bigg/ \left(\frac{R_0 n^2}{r} + 1\right)^2$$

Comparing this expression with (33) we see that the signal-to-noise ratio may be improved by increasing $R_0 n^2/r$, i.e., overcoupling the cavity until this noise source does not contribute any more. A comparison of (15) with (32) shows that overcoupling will not improve the signal-to-noise ratio when tuned to χ' . In this connection it should be pointed out that only those frequency stabilization schemes can alleviate the problem of frequency instabilities whose response time is at least of the order to the inverse modulation frequency since the troublesome noise components are at this frequency. Some stabilization schemes make use of the cavity into which the sample is placed as the stabilizing element. Although this system may be excellent for the observation of χ'' (it is the only one which can compensate for cavity microphonic), it fails in the case of χ' . The reason is that χ' makes itself observable essentially by a frequency shift which in this scheme would be compensated for.

* For a similar reason one cannot avoid an admixture of dispersion to an absorption signal, when investigating a saturated sample in which $\chi'_{max} \gg \chi''_{max}$.

C. Noise Due to Cavity Vibrations

A noise source which can be very troublesome at high modulation fields arises from the currents induced in the walls of the cavity from the modulating field. The interaction of these currents with the dc magnetic field causes mechanical vibrations of the cavity walls. This produces a signal when tuned to the dispersion, but to first order should give no signal when one is tuned to the cavity and sensitive to the absorption. However, any detuning will result in a signal, which, having the right frequency will pass through the narrow band amplifier and lock-in detector. Since this signal is proportional to the magnetic field, it will result in a background signal whose amplitude will vary as the magnetic field is being swept and thus causing a continuous shift in the base line. In a rectangular cavity this effect can be greatly reduced by a proper orientation of the cavity with respect to the dc magnetic field. This is due to the fact that by squeezing the broad face of a rectangular cavity (TE₁₀ mode) the frequency decreases, whereas by squeezing the narrow walls of the cavity the frequency increases. Thus in a proper orientation the two effects cancel each other out. We found another way of greatly reducing the effect by using a glass cavity having a silver coating* thick in comparison to a microwave skin depth but small in comparison to the modulation frequency skin depth, thereby decreasing the eddy currents without impairing the mechanical strength of the cavity.

D. Klystron Noise

There is very little data available on presently used klystrons. The data quoted by Hamilton, et al⁴ are on a 723A klystron. With an IF of 30 mc, bandwidth of 2.5-mc microwave output of 50 mw they obtained a noise power of 5×10^{-12} watts. Expressing their results in terms of a noise figure N_k such that the noise power output in the two side bands P_k is given by

$$P_k = 2N_k(kT\Delta\nu) \quad \therefore N_k = \frac{1}{2} \left(\frac{P_k}{P_0} \right) \frac{P_0}{kT\Delta\nu} = sP_0 \quad (34)$$

Substituting their numerical values one obtains for $s = 5000 \text{ Watt}^{-1}$. The values for s that we obtained with a 60 mc IF are:

Higher mode of V-153 klystron	$s \simeq 1,000 \text{ Watt}^{-1}$
Lower mode of V-153 klystron	$s \simeq 3,000 \text{ Watt}^{-1}$
Higher mode of X-13 klystron	$s \simeq 200 \text{ Watt}^{-1}$
Lower mode of X-13 klystron	$s \simeq 400 \text{ Watt}^{-1}$

* We are indebted to A. V. Hollenberg and V. J. DeLucca for the making of the glass cavities and to A. W. Treptow for the excellent silver coatings.

The method used to determine the above noise figures is similar to the one described in Reference 4. The figures are expected to be several times larger when a 30 mc IF is used. (A factor of 2 is quoted by Hamilton, et al.⁴) It is also worth noting that the relative noise power decreases on going to higher modes.

E. Signal-to-Noise Ratio for Specific Systems

In this section we will analyze specific systems under varying conditions. The reason why we do not present the analysis of one "The Best" system is that sometimes a compromise between complexity and sensitivity has to be reached and also because some systems may be superior at high power whereas others at low powers.

The expression for the noise power P_N at the output of the microwave detector (X-tal or bolometer) is⁵

$$P_N = (GN_K + F_{\text{AMPL}} + t - 1)(kT\Delta\nu) \quad (35)$$

where:

G = conversion gain of the detector (generally smaller than one. A quantity often used instead of G is the conversion loss $L = 1/G$).

N_K = noise figure at the input of the detector. Usually due to random amplitude or frequency fluctuations of the microwave source or the microwave components (see Section VID).

F_{AMPL} = noise figure of the amplifier

t = noise temperature of the detector

Comparing the equivalent voltage fluctuations of this noise power with the signal voltage as derived in (26), we get for the minimum detectable χ''

$$\chi_{\min}'' = \frac{1}{Q_0\eta\pi} \left[\frac{(GN_K + F_{\text{AMP}} + t - 1)kT\Delta\nu}{2GP_0} \right]^{\frac{1}{2}} \quad (36)$$

The above relation should apply to all systems. The problem then reduces to the determination of G , N_K , F_{AMP} , and t for the particular detection scheme. A difficulty arises from the fact that not only are those quantities a function of the RF power and modulation frequency but in the case of detectors vary from unit to unit. It is probably for this reason that the values quoted in the literature are sparse and are not in agreement with each other. The values used in this analysis for the X-band barretters (821) and crystals (1N23C) were obtained by us. Values for K-band crystals can be found in References 6 and 7. It should also be borne in mind that the values are time dependent and

will have to be modified as the "art" of detector manufacturing improves. Also new systems might come into prominence in the near future. An example would be the use of low noise travelling wave tubes preceding the detector or even low noise solid state masers.

1. Barretter (bolometer) detection.

The resistance of a barretter is given by the relation²

$$R = R_0 + kP^n \quad (37)$$

For practical purposes n may be taken as unity. The instantaneous power input to the barretter for a modulated microwave is given by

$$P = \frac{1}{R} \left[V_0 \sin \Omega t \left(1 + \frac{\Delta V}{V_0} \sin \omega t \right) + V_{dc} \right]^2 \quad (38)$$

where V_0 is the amplitude of the microwaves, ΔV the change of the amplitude due to the absorption given by (16), Ω the microwave frequency, ω the field modulation frequency, and V_{dc} the bias on the barretter. Expanding (38) and assuming that $\Delta V/V_0 \ll 1$ we get, after throwing out the high frequency terms,

$$R = R_0 + k \left(P_{RF} + P_{dc} + \frac{V_0 \Delta V}{R} \sin \omega t \right) \quad (39)$$

where P_{RF} is the power in the unmodulated carrier reaching the barretter. It is of course smaller than P_0 the microwave power from the klystron because of the power splitting in the magic T and the reflection from the cavity. Taking a reflection coefficient $\Gamma \simeq 0.5$ (see Section IVA), $P_{RF}/P_0 \simeq 0.1$. The desired voltage fluctuation associated with the resistance change is: $\delta V = I_0 \Delta R$, where

$$\Delta R = \frac{dR}{dP} \Delta P = \frac{dR}{dP} (\Delta P_{RF} + I_0^2 \Delta R), = \frac{dR}{dP} \left[\frac{1}{1 - I_0^2 \frac{dR}{dP}} \right] \Delta P_{RF} \quad (40)$$

$$\therefore \delta V \simeq I_0 k \frac{V_0 \Delta V}{(1 - I_0^2 k)} \sin \omega t$$

I_0 is the current bias on the barretter which we want to keep constant for a maximum voltage change δV .

The power gain of this device G is given by

$$\begin{aligned} G &= \frac{\text{Signal Power from Barretter}}{\text{Power in the Sidebands}} = \frac{\delta V^2}{2R} \bigg/ \frac{\Delta V^2}{4R} = 4 \frac{I_0^2 k^2 P_{RF}}{R(1 - I_0^2 k)} \\ &= 4k^2 \frac{P_{dc} P_{RF}}{R^2(1 - I_0^2 k)} \end{aligned} \quad (41)$$

The noise figure of the audio amplifier following the detector is given in general by:

$$F_{\text{AMP}} = \frac{R_{\text{equ}} + R_G}{R_G} \quad (42)$$

where R_G is the generator input resistance in this case the barretter resistance and R_{equ} is an equivalent noise resistor in series with the generators. The best input tube that we found was the General Electric GL 6072 triode* for which we measured an R_{equ} at 100 c.p.s. of $\sim 10^5 \Omega$. (This is due to flicker noise of the tube and has a $1/f$ dependence.) If we were to connect the barretter, having a resistance of $\sim 200 \Omega$ straight to the grid of the input tube we would get a noise figure of ~ 500 . However, by using a step-up transformer with a turns ratio $n > \frac{R_{\text{equ}}}{R_B}$ the noise figure of the amplifier can be reduced to nearly unity. We will assume in the following analysis that this has been done.

The noise temperature of the barretter t_B was thought to be approximately 2 since it is merely a platinum wire operating at an elevated temperature. To our surprise the measured value turned out to vary for different units between 4 and 40.† The noise figure was measured on about 20 different units obtained from 4 different manufacturers (P.R.D.; F.X.R. Narda, Sperry). The reason for this noise is not entirely clear at present. A possible explanation is the non-uniform heating of the wire which could set up air currents. They in turn can cool the wire in a random fashion giving rise to an additional noise component. An improvement of the noise figure was noted upon evacuating the barretter. The noise figure of a unit which was initially 10, dropped to the expected value of 2 after evacuation. However, it should be pointed out that this cannot be taken as a definite proof for the "air current theory" since the characteristics of the barretter changed markedly after evacuation. The sensitivity of the evacuated barretter went up from $5 \Omega/\text{mW}$ to $200 \Omega/\text{mW}$ which necessitated a reduction of the dc current from 8 to 1.5 mA. Also the response time went up by a factor of 20, so that the effectiveness of any noise mechanism with a $1/f$ spectrum would be greatly reduced. This approach however looks definitely promising in trying to design more sensitive and less noisy barretters. In the present work commercial unevacuated barretters were used, their noise temperature being taken as 4 in the following analysis. Under

* We are indebted to R. G. Shulman for bringing this tube to our attention.

† One unit which exhibited an extremely large noise figure of 1,000 was eliminated entirely. The solder point of the platinum wire was apparently defective.

those assumptions (36) becomes:

$$\chi''_{\text{MIN}} = \frac{1}{Q_0 \eta \pi} \left(\frac{kT \Delta \nu}{P_0} \right)^{\frac{1}{2}} \left(\frac{4 + N_k G}{G} \right)^{\frac{1}{2}} \quad (43)$$

a. Straight detection

A block diagram of the microwave part of a simple barretter system is shown in Fig. 6. The attenuator serves the purpose of preventing power saturation of the sample or burn out of the bolometer at high powers. By means of the slide screw tuner and magic T arrangement one makes the system sensitive to either the real or imaginary part of the susceptibility.

The characteristics of a typical barretter (like the Sperry No. 821) are: $R = 250 \, \Omega$; $k = 4.5 \, \Omega/mW$; $P_{\text{MAX}} = 32 \, mW$. We take the worst generator noise figure reported, i.e., $N_k = 5,000 P_{\text{RF}}$ (see Section VID).

The ratio of the minimum susceptibility $\chi''_{\text{MIN-OBS}}$ that can be detected with this system to the minimum theoretical value if one were limited by thermal noise only becomes with the aid of (41) and (43)

$$\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} = \left(\frac{4 + N_k G}{G} \right)^{\frac{1}{2}} = \left(\frac{1 + 5 \times 10^3 P_{\text{RF}} k^2 \frac{P_{\text{RF}} P_{\text{dc}}}{R^2 (1 - I_0^2 k)}}{k^2 \frac{P_{\text{RF}} P_{\text{dc}}}{R^2 (1 - I_0^2 k)}} \right)^{\frac{1}{2}} \quad (44)$$

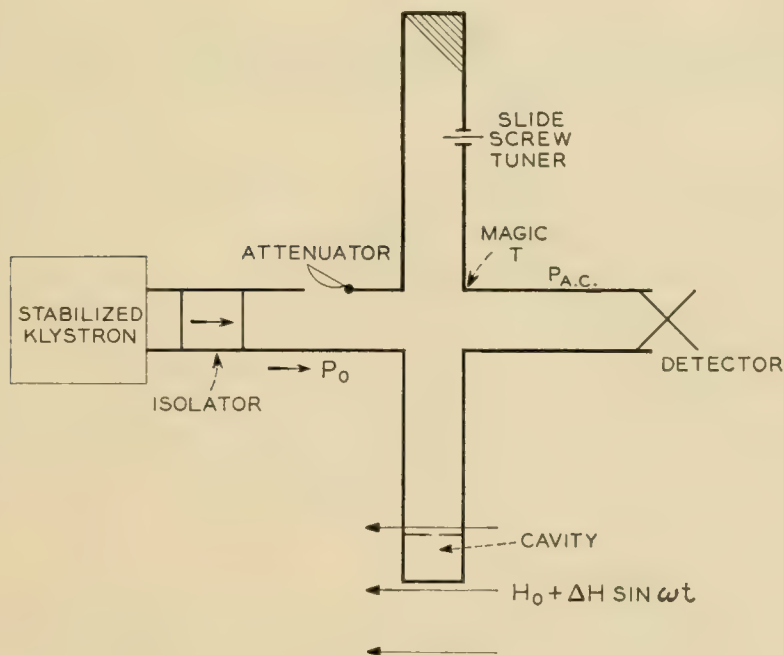


Fig. 6 — Essential microwave parts of a simple barretter or crystal set-up.

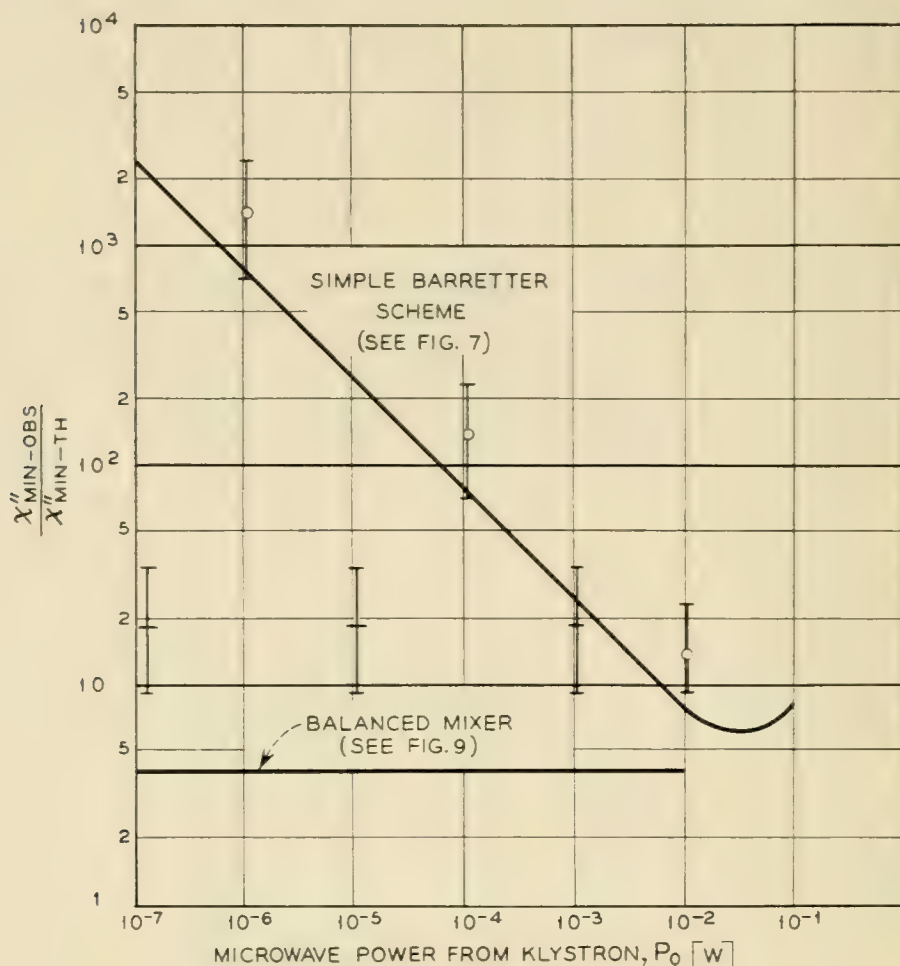


Fig. 7 — The ratio of the minimum detectable susceptibility to the minimum theoretical value versus microwave power for 2 different barretter schemes. Full lines correspond to the predicted sensitivity and dots indicate experimental values.

Equation (44) is plotted in Fig. 7. From this plot we see that the system is extremely poor at low powers (which is due to the low conversion gain of barretters) and also starts getting worse at high powers (due to the signal generator noise). The latter point is not of great importance since one can always buck down the microwave power by means of the slide screw tuner to the desired level. By using the evacuated barretter as mentioned earlier, the curve in Fig. 7 would be shifted to the left corresponding to the increased conversion gain.

b. *Balanced mixer detection*

An improved barretter scheme is shown in Fig. 8. It eliminates the poor conversion gain at low powers by employing a balanced mixer into which a large amount of microwave power P_2 can be fed from the same signal generator. Since the barretter noise should not be power

dependent (unlike in crystals) this procedure improves the conversion gain without increasing the noise. Since a balanced mixer is used the noise from the signal generator is also cancelled. A necessary precaution in this set-up is to include extra isolation between the second magic T and the mixer in order to prevent any microwave power from leaking through the balanced mixer into the cavity.

For this arrangement (44) becomes:

$$\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} \sqrt{2} \sqrt{\frac{4}{G}} \quad (45)$$

The factor of $\sqrt{2}$ arises from the fact that we had to split the power P_0 in the first magic T . Since in this scheme we are at liberty to vary the input power to the barretter we want to maximize G with respect to P_2 . For a fixed total power to the barretter given by its burn-out ratings (i.e., $P_2 + P_{\text{dc}} = \text{constant}$) (41) is a maximum for $P_2 \simeq P_{\text{dc}} \simeq \frac{P_{\text{MAX}}}{2}$. Taking again the data for the No. 821 barretter we get for $G_{\text{MAX}} \simeq 0.5$ and for

$$\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} \simeq 4 \quad (46)$$

Since the value of P_2 can be held constant irrespective of the power in the cavity, this ratio will be a constant (see Fig. 7).

It should be pointed out that in this system a wrong phasing of arm P_2 will result not only in a reduction of the signal, but also in an admixture of χ' and χ'' . Therefore after changing the power by means of a

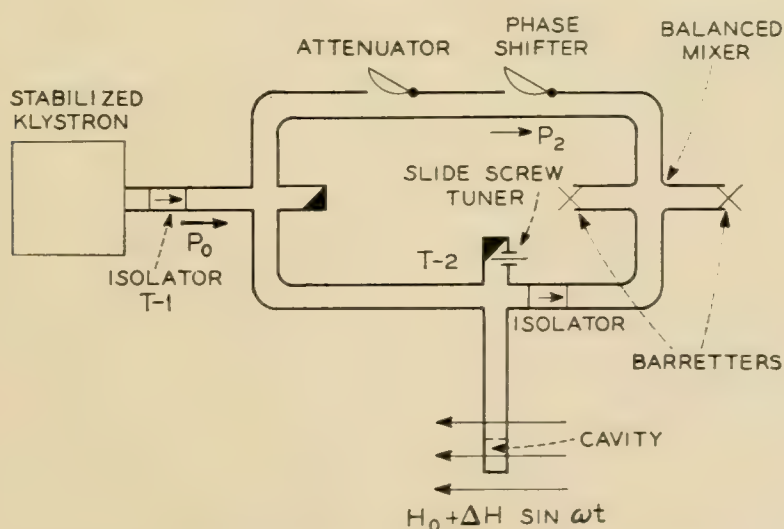


Fig. 8 — Barretter system with balanced mixer.

variable flap* attenuator (which also introduces a phase shift) or after changing the slide screw tuner the system has to be rephased again. This makes saturation measurements less convenient than in the superheterodyne system to be discussed in the next section. The obvious advantage of the homodyne detection scheme is that it requires only one microwave oscillator.

2. Crystal detection

A simple set-up is shown in Fig. 6. Although its microwave components are identical to the ones used in the barretter scheme, the analysis of this set-up is more complicated. The reason is that not only do crystal characteristics vary greatly from unit to unit but they cannot be described by one simple relation over the entire range of incident microwave power. One can roughly divide their characteristics into a square law region where the rectified current I is proportional to P_0 (holds for $P_0 < 10^{-5}$ Watts) and the linear region where I is proportional to $\sqrt{P_0}$. (holds for $P_0 > 10^{-4}$ Watts). The output noise of a crystal can be represented in general by the relation.^{5, 8, 9}

$$P_N = \left(\frac{\alpha I_0^2}{f} + 1 \right) kT\Delta\nu \quad \dagger \quad (47)$$

where f is the frequency around which the bandwidth $\Delta\nu$ is centered. This relation reduces for the square law region to:

$$P_N = \left(\frac{\beta P_{RF}^2}{f} + 1 \right) kT\Delta\nu \quad (48)$$

and for the linear region to

$$P_N = \left(\frac{\gamma P_{RF}}{f} + 1 \right) kT\Delta\nu \quad (49)$$

The average values of β we determined experimentally are:

$$\beta \simeq 5 \times 10^{14} \text{ Watt}^{-2} \text{ sec}^{-1} \quad \text{and}$$

$$\gamma \simeq 10^{11} \text{ Watt}^{-1} \text{ sec}^{-1}$$

The conversion gain G of the crystal can be represented by

$$G = SP_{RF} \quad (50)$$

in the square law region and by

$$G = \text{constant} = C \quad (51)$$

* The phase shift associated with the Hewlett-Packard X-382-A attenuator is quite small.

† Values of α for K-band crystals are quoted in References 6 and 7. They differ however from each other by approximately 3 orders of magnitude.

in the linear region. Values of S and C for the 1N23C were found to be $S \simeq 500 \text{ Watt}^{-1}$ and $C \simeq 0.3$.

a. Simple straight detection

If one does not make use of the bucking possibilities of the magic T (i.e., eliminate the slide screw tuner in Fig. 6) one has the simplest possible set-up sensitive to χ'' . Under those conditions the microwave power reaching the crystal will be identical to the reflected power from the cavity. Equation (36) becomes:

$$\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} = \left(\frac{GN_k + F_{\text{AMP}} + t - 1}{G} \right)^{\frac{1}{2}} \quad (52)$$

With the aid of (48), (49), (50), and (51), this relation reduces for the 1N23C in the square law region to:

$$\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} = \left(\frac{1 + 5 \times 10^9 P_0^2}{50 P_0} \right)^{\frac{1}{2}} \quad (53)$$

and for the linear region to:

$$\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} = (3 \times 10^7 P_0)^{\frac{1}{2}} \quad (54)$$

A plot of (53) and (54) is shown in Fig. 9. As before the assumption was made that $P_{\text{RF}}/P_0 \simeq 0.1$ (see barretter case). The noise figure of the amplifier F_{AMP} was taken as unity which again can be closely approached by means of a step-up transformer. The field modulation frequency was assumed to be 1,000 c.p.sec., although (47) shows that from a point of view of noise one would like to go to as high a frequency as possible. However practical consideration such as power requirements for getting a given modulation field, pick-up problems, skin depth losses in the cavity wall usually set an upper limit. The modulation frequency may be also dictated at times by the relaxation times of the investigated sample.¹⁰

b. Straight detection with optimum microwave bucking

From Fig. 9, we see that the straight crystal detection scheme suffers at low powers because of the poor conversion gain of the crystal and at high powers because of excess crystal noise. This situation can be greatly improved by adding some microwave power to the crystal when the reflected power from the cavity is low (to be referred to as positive bucking) or subtracting some of the power in the other case (negative bucking). In this section we will find the improvement over the unbucked system and the amount of bucking required to effect it.

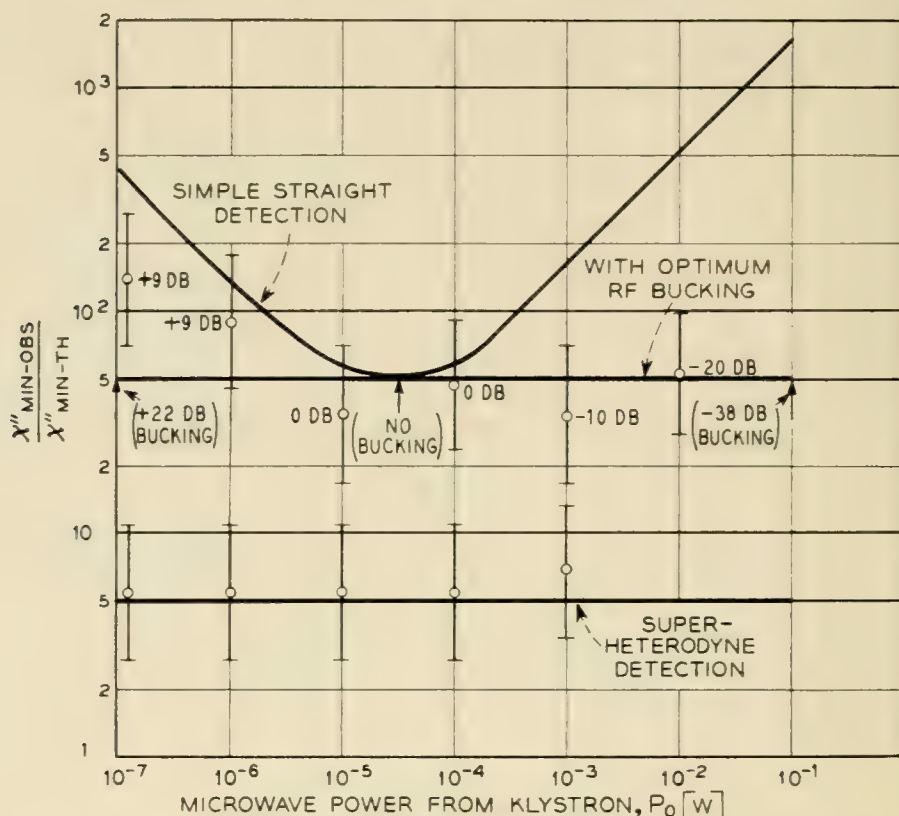


Fig. 9 — The ratio of the minimum observable susceptibility to the minimum theoretical value versus microwave power for different crystal detection schemes. Full lines correspond to the predicted sensitivity and dots indicate experimental values.

We define the bucking parameter B by the relation

$$P_x = BP_{RF} \quad (55)$$

Where P_{RF} is the microwave power at the crystal before and P_x after the bucking is applied. We further assume that after the bucking is applied the crystals will operate in the square low region. Combining (49), (50), and (52) and neglecting the term GN_k which is small in comparison to the other term we get for the bucking scheme:

$$\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} = \left(\frac{F_{\text{AMP}} + \frac{\beta B^2 P_{RF}^2}{f}}{SBP_{RF}} \right)^{\frac{1}{2}} \quad (56)$$

In order to find the optimum bucking parameter, we set

$$\frac{d}{dB} \left(\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} \right) = 0 \text{ which results in} \quad (57)$$

$$\left(B = \frac{F_{\text{AMP}} f}{\beta P_{RF}^2} \right)^{\frac{1}{2}}$$

Putting in the numerical values for the $1N23C$ as quoted previously we get that $B = 1.4 \times 10^{-6}/P_{RF}$. Equation (56) becomes with the optimum bucking parameter

$$\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} = \left[\frac{2F_{\text{AMP}}}{S \left(\frac{F_{\text{AMP}} f}{\beta} \right)^{\frac{1}{2}}} \right] \quad (58)$$

For the case under discussion this ratio turns out to be $\simeq 50$ independent of P_0 . Fig. 9 shows a plot of (58). The values in parenthesis indicate the degree of bucking necessary to accomplish this ratio as determined from (57). It should be noted that the negative bucking can be easily accomplished by means of the slide screw tuner in the magic T arm (see Fig. 6) whereas for large positive buckings a scheme like in Fig. 8 has to be used. (Some positive bucking can of course be also accomplished by means of the slide screw tuner).

c. The superheterodyne scheme

The RF bucking system just described bears a certain resemblance to the balanced mixer barretter scheme. In both cases additional microwave power was added to the detector in order to increase the conversion gain. However in the crystal scheme this resulted in an increase in noise power whereas this should not be the case with barretters. The question arises whether a decent conversion gain in crystals has to be always accompanied by a large noise power. An inspection of (49) shows that around frequencies of tens of megacycles* or higher the noise output of the crystal becomes negligible. As pointed out earlier such high magnetic field modulation frequencies are not feasible. However in a superheterodyne system the crystal outputs will be at an intermediate frequency of 30 or 60 mc. This will make the flicker noise components negligible even at high powers where the conversion gain is good. The conventional way to obtain the intermediate frequency is to beat the reflected signal from the cavity with a local oscillator (see Fig. 10) which is removed from the signal generator by the I.F. frequency. In order to eliminate the noise from the local oscillator a balanced mixer should be employed. The ratio

$$\left(\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} \right) \text{ becomes then from equ. 52 } \left(\frac{F_{IF} + t - 1}{G} \right)^{\frac{1}{2}} \quad (59)$$

The expression in the brackets is called in radar work¹¹ the overall

* It was shown by G. R. Nicoll⁹ that this equation holds up to this frequency range.

noise figure of the receiver F . We found that a noise figure of about 11–14 db is easily attainable with commercial I.F. amplifiers and balanced mixer. This would give us a ratio of

$$\frac{\chi''_{\text{MIN-OBS}}}{\chi''_{\text{MIN-TH}}} \simeq 5$$

which is plotted together with the other crystal schemes in Fig. 9. Although this system does necessitate 2 stable microwave sources, it is not difficult to operate once they are set-up. This was not considered as a major disadvantage at least not at X-band. The phasing problem discussed in connection with the mixer barretter scheme of comparable sensitivity is eliminated. An additional small advantage is the ruggedness of crystals in comparison to barretters and the availability of good commercial balanced crystal mixers. There are other double frequency schemes which do not need 2 separate microwave signal generators. The other frequency may be obtained by amplitude or phase modulating one signal generator by an IF frequency. The side bands which are produced in this way are displaced by just the IF frequency and may be utilized instead of the second signal generator. Schemes of this sort look particularly promising for frequencies well above X-band in which case it might prove difficult to maintain the difference frequency of two separate microwave generators within the band width of the IF.

F. Experimental Determination of Sensitivity Limits

1. Preparation of samples

In order to get an experimental check on the previous analysis, samples with a known number of spins had to be prepared. Two sets of samples were made. One consisted of single $\text{CuSO}_4 \cdot 5\text{H}_2\text{O}$ crystals of varying sizes hermetically sealed between 2 sheets of polyethylene. The other set consisted of different amounts of diphenyl piceryl hydrazyl* which were similarly sealed up. D.P.H. samples having less than 10^{17} spins were prepared by dissolving known amounts of the free radical in benzene and putting a drop of this solution on the polyethylene. After the benzene had evaporated, it was sealed up with another sheet of polyethylene. The g -values of $\text{CuSO}_4 \cdot 5\text{H}_2\text{O}$ and D.P.H. differ enough so that both samples can be conveniently run simultaneously. This was done in order to check the self consistency of the two sets of samples. The measured integrated susceptibility of all the D.P.H. samples

* We are indebted to A. N. Holden for supplying us with this material.

with more than 10^{16} spins agreed within a few percent with the calculated value. The calculated value being based on the known amount of D.P.H. and the measured value being referred to the known amount of $\text{CuSO}_4 \cdot 5\text{H}_2\text{O}$. D.P.H. samples with less than 10^{16} spins had all a smaller number of effective spins than calculated. The discrepancy was more pronounced the smaller the sample. There was also evidence that the smaller D.P.H. samples deteriorated with time. As a typical example we quote a sample which started out as 10^{15} effective spins and was reduced after 4 weeks to 4×10^{14} effective spins and another one which initially had 10^{14} spins, deteriorated in the same time interval to 10^{13} spins. Since only the smaller samples were noticeably affected, this deterioration seems to be associated with a surface reaction. It was also observed that the line width between inflection points of the D.P.H. samples with less than 10^{15} spins increased from 1.8 oersteds to 2.7 oersteds. This broadening probably arises from a reduction in the exchange narrowing mechanism due to the spreading out of the sample.

2. Comparison of experimental results with theory

In checking the sensitivity of the equipment D.P.H. samples were used and the signal to noise was estimated from the recorded output. The experimental points thus obtained are shown in Fig. 7 and Fig. 9. We believe that the results are significant to within a factor of 2, the main error arising from the estimate of the RMS noise. The band width of the lock-in detector was $\Delta\nu = 0.03 \text{ sec}^{-1}$, $Q_0 = 4,000$, and the field modulation used was 3 oersteds p.t.p., 100 c.p.sec. for the barretter schemes and 1,000 c.p.sec. for the crystal schemes. This large modulation field somewhat distorts the line, but, as mentioned earlier was done in order to get the full signal. The D.P.H. samples were calibrated against $\text{CuSO}_4 \cdot 5\text{H}_2\text{O}$ before each run. Even so it was not felt safe to use samples which had less than 10^{13} spins.

Referring to Fig. 8 we see that for the straight barretter detector the experimental points agree fairly well with the predicted value, but in the balanced mixer scheme fall short by about a factor of 4. A possible explanation of this discrepancy is that the barretters were not completely matched in which case the noise from the local oscillator would not be compensated for.

Fig. 9 shows the experimental points for the crystal schemes. For powers between 10^{-7} W and 10^{-5} W the system used fell between the simple straight detection scheme and the one utilizing optimum RF bucking. The reason is that it was very easy to obtain a certain amount of positive bucking (+9 db) by merely adjusting one arm of the magic T .

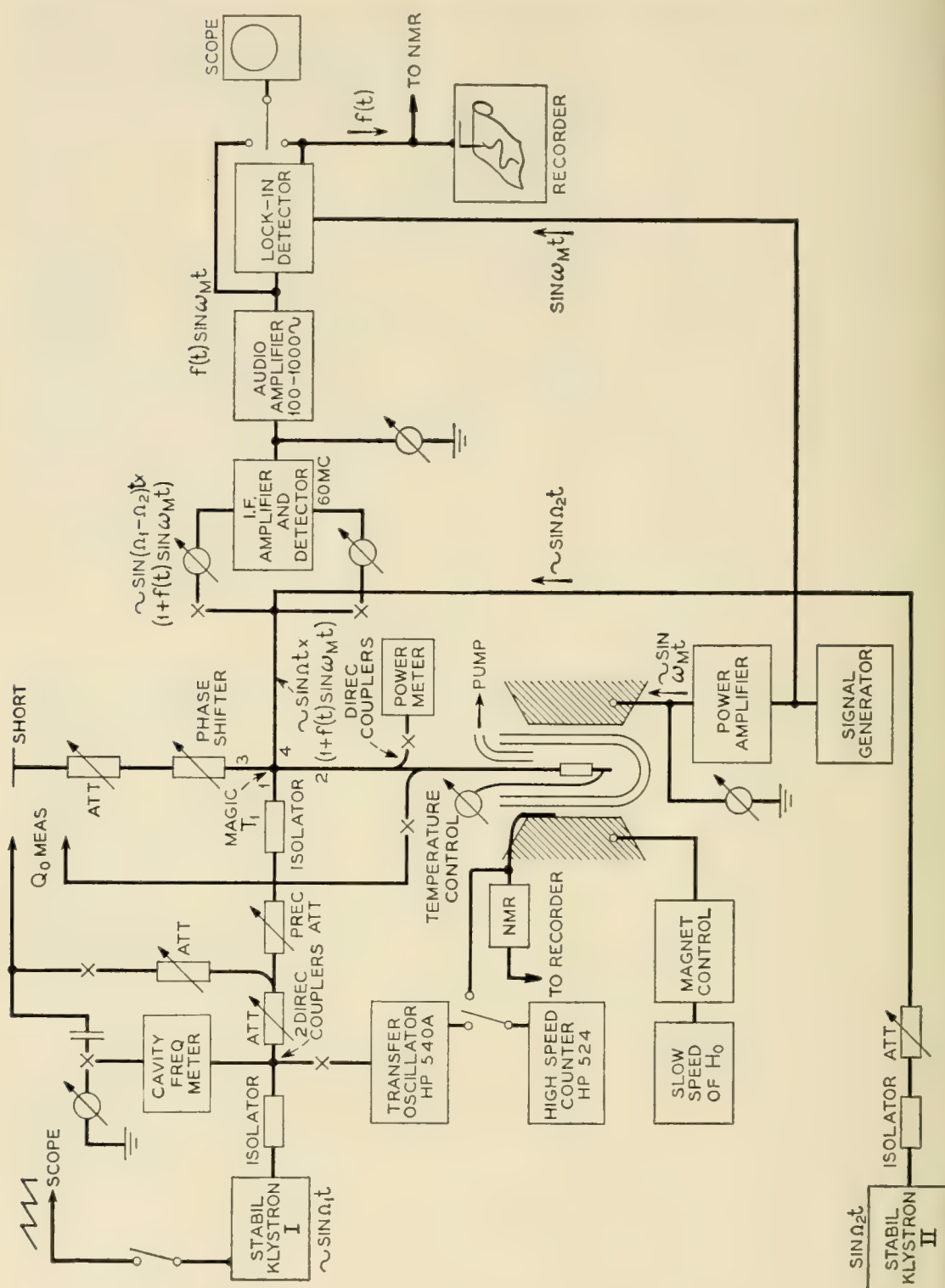


Fig. 10 — Block diagram of a superheterodyne paramagnetic resonance spectrometer

It would however have been a great deal more difficult to obtain the entire bucking of +22 db at 10^{-7} W since a set-up like in Fig. 8 would have to be used. Thus for the sake of simplicity the extra factor in signal to noise of 2 or 3 was abandoned. The amount of negative bucking at the higher powers will be limited by the stability of the bridge. A practical limit of (40–50) db was characteristic of our set-up. We see from Fig. 9 that the agreement between the experimentally determined sensitivity and the theoretically predicted sensitivity is satisfactory.

The experimental results on the superheterodyne scheme agrees again very well with the predicted values up to a power level of 10^{-3} W. (This corresponds to less than 10^{12} spins in D.P.H.) Above this level T_1 (see fig. 10) has to be balanced to better than 40 db to keep the IF carrier amplitude within the required value. Instabilities in the bridge due to mechanical vibrations and thermal drifts start to contribute to the noise. Thus at high power levels the superhet scheme starts to lose some of its advantages unless special precautions are being taken to eliminate the above mentioned noise factors. A great deal in this direction could probably be accomplished by shock-mounting the microwave components and better temperature stability for slow drifts. Since we were mainly interested in powers below 1 mW, our efforts were limited to controlling the temperature of the room to $\pm 1^\circ\text{C}$.

Since the superhet scheme was found to be the most sensitive one, it might be worthwhile to discuss it in more detail. A block diagram of the set up is shown in Fig. 10.

The signal generator feeds into the magic T , where its power is split between arm 2 and 3. Arm 2 has the reflection cavity with the sample, the reflected voltage being bucked out with the aid of arm 3. For this purpose arm 3 has a phase shifter and attenuator, an arrangement which was found to be more satisfactory than a slide screw tuner as far as stability and ease of operation goes. The desired signal appears then in arm 4. It is fed into a balanced mixer which receives the local oscillator power from the stabilized klystron II. The output of the balanced mixer is then fed through the IF amplifier, detector, audio amplifier and lock-in detector. The circuits of each of those components is fairly standard and will not be dwelled upon further. The microwave power is measured in arm 2 of the magic T . The power reflected from the cavity is also monitored in arm 2. This is of great help in finding the cavity when klystron I is swept in frequency by means of a sawtooth voltage on its reflector. Since the klystron mode itself might have some dips in it, (which might be mistaken for the cavity), it proved helpful to display on the scope the klystron mode simultaneously with the reflected power

from the cavity. This also provides a convenient way to measure the Q_0 of the cavity.¹² The frequency is measured roughly by means of a cavity frequency meter and more precisely by means of a transfer oscillator and high speed counter. The magnetic field is measured by means of a nuclear magnetic resonance set-up, its frequency being measured on the same counter as the microwave frequency. The nuclear resonance signal is recorded on the same trace as the electron resonance signal. Thus if the magnetic field is homogeneous enough, the nuclear sample will see the same field as the electronic sample and g -values can be conveniently determined to the accuracy of the nuclear moment (this also assumes that the signal is large enough, so that no additional error is introduced in determining the exact location of the resonance.) The field modulation coils are mounted on the pole faces and are energized by a 50-watt power amplifier. A field of 50 oersteds p.t.p. is available at 1,000 cps and a slightly higher field at 100 cps.

The magnet is a Verian 12" modified so that it can rotate around an axis perpendicular to H_0 . This was done mainly in order to make anisotropy measurements more convenient. This enables one to make quick saturation measurements in isotropic materials without having to change the incident RF power. This is accomplished by rotating the magnetic field and measuring the signal strength versus angle. Since only the RF field perpendicular to the dc field causes transitions, the signal in an unsaturated isotropic sample should go as $\cos^2 \theta$; where θ is the angle between H_1 and H_0 . From the deviation from this dependence, the saturation parameter can be found. This could also be done by rotating the cavity, but at microwaves is not as easy as rotating the field.

VII. A NOTE ON THE EFFECTIVE BANDWIDTH

There seems to be some confusion as to how narrow one should make an audio amplifier preceding a phase sensitive detector (lock-in) or why the band width of the IF amplifier doesn't enter in a superhet scheme. Those and similar questions have to do with the effective band width of the system $\Delta\nu$ which appears in (26). Since similar questions have been rigorously analyzed by other authors,^{13, 14} the present discussion will try to stress some of the physical ideas underlying the different detection schemes.

We consider first the simple scheme illustrated in Fig. 11. It consists of an amplifier with band width $\Delta\nu_1$ centered around ν_1 followed by a phase sensitive detector with a reference voltage at V_1 . The output of the phase sensitive detector has an RC filter of band width $\Delta\nu_2$. One can see that in such a system the only noise components centered around

ν_1 (this being also the reference frequency) in a band width $\Delta\nu_2$ will contribute to the output noise. This is because the beat between 2 noise components like ν_n and ν_m (see Fig. 11) is too far removed from ν_1 to produce an output voltage. (This statement implies the condition that $\Delta\nu_1 < \nu_1$ otherwise the beat between $2\nu_1$ and ν_1 could come through.) Thus in this system the band width of the amplifier is immaterial as long as the noise voltages are not so large as to saturate it.

A more serious situation may arise in the absence of a reference voltage. In this case the noise components within the band width $\Delta\nu_1$ can beat with each other and produce a noise output which would increase with the band width. This could become especially detrimental in a superheterodyne scheme in which the IF band width can be a million times larger than the output band width. It can be shown, however, that if the carrier voltage V_c at the output of the IF is large enough the IF bandwidth ΔF_{IF} does not enter into the noise consideration¹³ the criterion essentially is that

$$V_c^2 > G^2 k T Z \Delta F_{IF} \quad (60)$$

where G is the IF amplifier gain, and Z the input impedance. Condition (60) means that we want the noise which beats with the carrier to be greater than the beat between 2 noise terms. Since the former is proportional to the carrier, its predominance can be easily ascertained experimentally by increasing the IF carrier and noting whether the noise output increases proportionally. If it does, (60) is fulfilled.

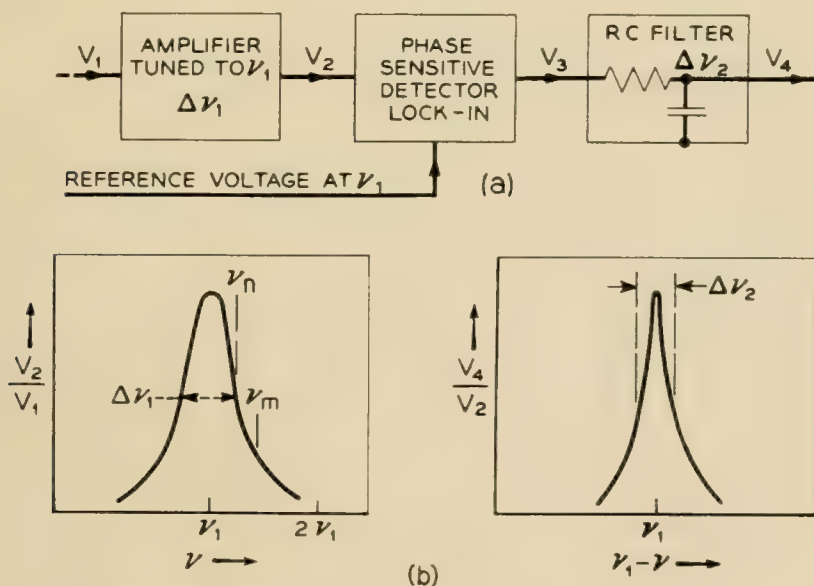


Fig. 11 — Effective band width of a phase sensitive detector $\Delta\nu_{eff} = \Delta\nu_2$. Note that the band width of the amplifier does not enter as long as $\Delta\nu_1 < \nu_1$.

In order to see what maximum gain G (60) imposes on a typical system we assume $\Delta F_{\text{IF}} = 5 \times 10^6$ c.p.sec. $Z = 10^3 \Omega$; $V_c \simeq 1\text{V}$. Under those conditions we get from (60) that G has to be smaller than approximately 10^5 . If on the other hand G is very small the signal level at the audio amplifier input is so low that the flicker noise of the detector can still come in. A good practical figure for the IF amplifier gain is around 60 db.

VIII. SATURATION EFFECTS

In all the previous considerations RF power saturation effects were neglected, i.e., we have assumed that the power absorbed is proportional to H_1^2 , where H_1 is the RF magnetic field. When this assumption is no longer satisfied, the question of sensitivity has to be re-examined for different degrees of saturation. However it is difficult from an experimental point of view to change the conditions of the experiment for each degree of saturation and therefore an elaborate analysis of this case does not seem to be warranted. However it might be of interest to see the effect on the in phase component of the signal at complete or nearly complete saturation

The change in output voltage for a reflection cavity is (15)

$$\Delta V = \sqrt{2} V \frac{R_0 n^2}{(R_0 n^2 + r)^2} \Delta r$$

and from (4)

$$\frac{\Delta r}{r} = \frac{\Delta Q}{Q_0} = 4\pi\eta Q_0 \chi''$$

The RF magnetic field in the cavity is given by

$$H_1^2 = C Q_0 (1 - \Gamma^2) P_{in} \quad (71)$$

where C is a constant dependent on the geometry of the cavity the reflection coefficient. Assuming a simple homogeneous saturation behaviour we substitute for χ'' the saturated value of $\chi_s''^{(1)}$.

$$\chi_s'' = \frac{\chi_u''}{1 + \gamma_1^2 H_1^2 T_1 T_2} = \frac{\chi_u''}{1 + \gamma_1^2 T_1 T_2 C Q_0 (1 - \Gamma^2) P_{in}} \quad (72)$$

where χ_u'' is the unsaturated value of the susceptibility and P_{in} the power from the microwave source.

$$\frac{\Delta r}{r} = 4\pi\eta Q_0 \frac{\chi_u''}{1 + \gamma_1^2 T_1 T_2 Q_0 (1 - \Gamma^2) P_{in}} \quad (73)$$

and the output voltage ΔV .

$$\Delta V = \sqrt{2} V_0 \frac{R_0 n^2}{(R_0 n^2 + r)^2} r^4 \pi \eta Q_0 \frac{\chi_u''}{1 + \gamma_1^2 T_1 T_2 C Q_0 (1 - \Gamma^2) P_{in}} \quad (74)$$

For a high degree of saturation

$$\gamma_1^2 T_1 T_2 C Q_0 (1 - \Gamma^2) P_{in} \gg 1$$

and substituting for

$$\Gamma = \frac{\frac{R_0 n^2}{r} - 1}{\frac{R_0 n^2}{r} + 1}$$

we get:

$$\Delta V = \frac{\sqrt{2} V_0 \pi \eta \chi_u''}{P_{in} \gamma_1^2 T_1 T_2 C} \quad (75)$$

The above relation shows that under saturated conditions the Q of the cavity does not enter and one might as well not use one or use a very much overcoupled cavity. This is one of the reasons why in microwave gas spectroscopy,* where lines are easier saturated a cavity is not used. (The more important reason is that in most cases one sweeps the frequency of the source, so that a cavity is difficult to use.)

Equation (75) also shows that the signal and also signal to noise goes down with increasing RF power. The above argument does not hold for the out-of-phase (dispersion) signal, in particular it breaks down completely for signals observed under fast adiabatic passage conditions.¹ For the latter case one wants as high an RF field as possible.

IX. ACKNOWLEDGEMENT

I profited greatly from discussions with various members of the resonance group at the University of California in particular with Profs. A. F. Kip, and A. M. Portis and at Bell Telephone Laboratories with Drs. R. C. Fletcher and S. Geschwind. I would like especially to thank E. Gere for his expert help in the construction of the equipment and to Prof. C. P. Slichter and Dr. R. H. Silsbee for helpful criticism of the manuscript.

REFERENCES

1. See for example F. Bloch, Phys. Rev., **70**, p. 460, 1946. N. Bloembergen, E. M. Purcell, R. V. Pound, Phys. Rev., **73**, p. 679, 1948.

* In microwave gas spectroscopy the fractional power loss per unit length α is used. Its relation to the susceptibility is $\alpha = 8\pi^2 \chi'' / \lambda_g$ where λ_g is the guide wavelength.

2. Montgomery, Technique of Microwave Measurements, Rad. Lab. Series, No. 11.
3. C. H. Townes and S. Geschwind, J.A.P., **19**, p. 795, Aug., 1948.
4. Hamilton, Knipp and Kupper, Klystrons and Microwave Triodes. McGraw Hill, 1948, Rad. Lab. Series No. 6, p. 475.
5. Torrey, H. C. and Whitmer, C. A., Crystal Rectifiers, Rad. Lab. Series, Vol. 15, McGraw Hill, 1947.
6. M. W. P. Strandberg, H. R. Johnson, J. R. Eshbach, R.S.I., **25**, pp. 776-792, Aug., 1954.
7. Townes and Schawlow, Microwave Spectroscopy, McGraw Hill, 1955.
8. Miller, P. H., Noise Spectrum of Crystal Rectifiers, Proc. I.R.E., **35**, p. 252, 1947.
9. G. R. Nicoll, Noise in Silicon Microwave Diodes, Proc. I.E.E., **101**, pp. 317-29, Sept., 1954.
10. See for example K. Holbach, Helv. Physica Acta, **27**, p. 259, 1954; A. M. Portis, Phys. Rev., **100**, p. 1219, 1955.
11. Pound, R. V., Microwave Mixers, M.I.T. Radiation Lab. Series 16, McGraw Hill.
12. E. D. Reed, Proceedings of the National Electronics Conference, **7**, p. 162, 1951.
13. S. O. Rice, B.S.T.J., **23**, pp. 282-332; B.S.T.J., **24**, pp. 46-156, 1945.
14. A. van der Ziel, Noise, Prentice Hall, Inc., 1954.

The Determination of Pressure Coefficients of Capacitance for Certain Geometries

By D. W. McCALL

(Manuscript received February 15, 1955)

Expressions are derived for the pressure coefficients of capacitance of parallel plate capacitors subjected to one-dimensional and hydrostatic pressures and of cylindrical capacitors subjected to radial compression. The derivations apply to systems in which the dielectrics are isotropic, elastic solids.

I. INTRODUCTION

The electrical capacitance between two conductors separated by a dielectric is a quantity which can be calculated with ease only in certain geometrical arrangements of high symmetry. Even the classic example of parallel plates presents major difficulties as one may only perform the calculation exactly for the case of plates of infinite area or vanishing separation. The approximation becomes poor when $(\text{area})^{1/2}/(\text{separation})$ becomes small and the theoretical treatment of edge effects is sufficiently difficult that it has not been solved though the solution would greatly facilitate dielectric constant measurement.

When pressure enters into the situation as a variable the difficulties are enhanced as one must be able to describe the geometry effects as well as the change in dielectric constant.

The engineers responsible for designing submarine cables are confronted with the necessity of knowing the manner in which capacitance depends upon pressure as may be illustrated in the following way. A submarine telephone cable is composed of a central copper conductor surrounded by a sheath of dielectric material. Due to the extreme length repeaters must be placed at intervals, the separation being determined by the attenuation of the cable. The attenuation, α , of a coaxial telephone cable may be written

$$\alpha = (G/2)(L/C)^{\frac{1}{2}} + (R/2)(C/L)^{\frac{1}{2}}$$

where G is the conductance of the dielectric per unit length, C the

capacitance, L the inductance, and R the conductor resistance per unit length. The second term contributes about 99% of the attenuation. Considering only this term we deduce

$$(1/\alpha)(\partial\alpha/\partial P) \cong (1/R)(\partial R/\partial P) + (1/2C)(\partial C/\partial P) - (1/2L)(\partial L/\partial P)$$

Accurate knowledge of the coefficient $(1/C)(\partial C/\partial P)$ is thus essential in designing very long cables which are to be exposed to high pressures.

In evaluating dielectric materials for use in cables it is often desirable to make measurements on sheet specimens rather than cable. It thus becomes necessary to be able to translate sheet data into cable data. It is the purpose of this paper to analyze the problem of calculating pressure coefficients of capacitance for certain simple geometries and to consider the methods of measurement which have been used. It will be shown that results of theory and experiment are in as good agreement as can be expected but more accurate measurements of electric and elastic properties are needed.

The equations which will be derived are also necessary if one wishes to determine the dependence of dielectric constant on pressure using any of the geometries described herein.

The problems treated in this paper are particularly simple and amenable to mathematical treatment but many problems encountered in submarine cable design are at present subject to solution only by empirical means. Fundamental investigations of the effects of pressure on dielectric materials are needed.

II. THEORETICAL TREATMENT

In the following treatment we consider that the dielectric substance is an elastic solid which obeys Hooke's law. We denote the relative permittivity or dielectric constant by ϵ , the permittivity of free space by ϵ_0 ,* the principal stresses and strains by τ_{ii} and e_{ii} , the density by ρ , the compressibility by k , and Poisson's ratio by σ .

A. Calculation of $\frac{1}{\epsilon} \frac{\partial \epsilon}{\partial P}$

One of the quantities which will be needed in the evaluation of

$$\frac{1}{C} \frac{\partial C}{\partial P} \quad \text{is} \quad \frac{1}{\epsilon} \frac{\partial \epsilon}{\partial P}$$

As ϵ is not dependent on the geometric configuration it can be calculated

* $\epsilon_0 = 8.86 \times 10^{-12}$ farads/meter.

in general and the result applied to each of the special cases to follow. A relation between dielectric constant and density is required and usually, when dealing with non-polar dielectrics, one assumes that the Clausius-Mosotti relation gives the proper dependence. That is

$$\frac{\epsilon - 1}{\epsilon + 2} = (\text{constant}) \rho \quad (1)$$

This formula may be differentiated to give

$$\frac{1}{\epsilon} \frac{\partial \epsilon}{\partial P} = \frac{(\epsilon - 1)(\epsilon + 2)}{3\epsilon} \frac{1}{\rho} \frac{\partial \rho}{\partial P} \quad (2)$$

In the theory presented herein, (1) will be used though it is at best an approximation. Corrections to the Clausius-Mosotti formula¹ which have been given do not seem applicable to polymer dielectrics and introduce parameters which must be fitted.

B. The effect of a One-Dimensional Pressure Acting on a Disc

Consider a one-dimensional pressure, $-P$, acting along the axis of a circular disc of dielectric material with electrodes affixed to opposite faces. Assume the disc is constrained such that no lateral displacement can occur. Let t be the thickness and A the area of the disc.

The capacitance of such a capacitor is given by the equation

$$C = \epsilon \epsilon_0 \frac{A}{t}$$

so the desired pressure coefficient is

$$\frac{1}{C} \frac{\partial C}{\partial P} = \frac{1}{\epsilon} \frac{\partial \epsilon}{\partial P} - \frac{1}{t} \frac{\partial t}{\partial P} \quad (3)$$

where use has been made of the condition that the area is constant (i.e., no lateral displacement). Hooke's law states

$$e_{xx} = \frac{k}{3(1 - 2\sigma)} [\tau_{xx} - \sigma(\tau_{yy} + \tau_{zz})] \quad (4)$$

$$e_{yy} = \frac{k}{3(1 - 2\sigma)} [\tau_{yy} - \sigma(\tau_{xx} + \tau_{zz})] \quad (5)$$

$$e_{zz} = \frac{k}{3(1 - 2\sigma)} [\tau_{zz} - \sigma(\tau_{xx} + \tau_{yy})] \quad (6)$$

¹ C. J. F. Böttcher, *Theory of Electric Polarisation*, Elsevier Publishing Co., Amsterdam, 1952, p. 199 et. seq.

We assume the z -axis lies along the disc axis so $\tau_{zz} = -P$ and by symmetry $\tau_{xx} = \tau_{yy}$. The condition of no lateral strain states $e_{xx} = e_{yy} = 0$ which when combined with (4) gives

$$\tau_{xx} = \tau_{yy} = -\frac{\sigma}{1-\sigma} P$$

Using the last result with (6) we obtain

$$e_{zz} = -\frac{kP}{3} \left(\frac{1+\sigma}{1-\sigma} \right)$$

and

$$\frac{e_{zz}}{P} = \frac{1}{t} \frac{\partial t}{\partial P} = -\frac{1}{\rho} \frac{\partial \rho}{\partial P} = -\frac{k}{3} \left(\frac{1+\sigma}{1-\sigma} \right) \quad (7)$$

Combination of (2), (3), and (7) results in

$$\frac{1}{C} \frac{\partial C}{\partial P} = \left[\frac{(\varepsilon-1)(\varepsilon+2)}{3\varepsilon} + 1 \right] \frac{k}{3} \left(\frac{1+\sigma}{1-\sigma} \right) \quad (8)$$

C. The Effect of a Hydrostatic Pressure Acting on a Disc

Assume that conditions are similar to those considered in section B except that $e_{xx} = e_{yy} = 0$ is now replaced by

$$\tau_{xx} = \tau_{yy} = \tau_{zz} = -P.$$

The area is no longer independent of pressure so

$$\frac{1}{C} \frac{\partial C}{\partial P} = \frac{1}{\varepsilon} \frac{\partial \varepsilon}{\partial P} - \frac{1}{t} \frac{\partial t}{\partial P} + \frac{1}{A} \frac{\partial A}{\partial P} \quad (9)$$

Hooke's law, (4), (5), (6), now becomes

$$e_{ii} = \frac{-kP}{3} \quad (i = x, y, z)$$

Thus

$$\frac{1}{\rho} \frac{\partial \rho}{\partial P} = -\frac{(e_{xx} + e_{yy} + e_{zz})}{P} = k \quad (10)$$

$$\frac{1}{A} \frac{\partial A}{\partial P} = \frac{(e_{xx} + e_{yy})}{P} = -\frac{2k}{3} \quad (11)$$

$$\frac{1}{t} \frac{\partial t}{\partial P} = \frac{e_{zz}}{P} = -\frac{k}{3} \quad (12)$$

Combination of (2), (9), (10), (11), and (12) results in

$$\frac{1}{C} \frac{\partial C}{\partial P} = \left[\frac{(\varepsilon - 1)(\varepsilon + 2)}{3\varepsilon} - \frac{1}{3} \right] k \quad (13)$$

D. The Effect of a Radial Pressure Acting on a Cylindrical Annulus

Consider a radial pressure acting normally to the axis of a cylindrical annulus to which electrodes are affixed to the inner and outer surfaces. Let the inner and outer radii be a and b respectively and assume the cylinder is filled with an incompressible substance so that the inner radius is not pressure dependent. The capacitance per unit length is given by

$$C_l = \frac{2\pi\varepsilon\varepsilon_0}{\ln \frac{b}{a}}$$

so

$$\frac{1}{C} \frac{\partial C}{\partial P} = \frac{1}{\varepsilon} \frac{\partial \varepsilon}{\partial P} - \frac{1}{\ln \frac{b}{a}} \frac{1}{b} \frac{\partial b}{\partial P} \quad (14)$$

We employ cylindrical coordinates (r, θ, z) where the z -axis is taken along the cylinder axis. Hooke's law becomes

$$e_{rr} = \frac{k}{3(1 - 2\sigma)} [\tau_{rr} - \sigma(\tau_{\theta\theta} + \tau_{zz})] \quad (15)$$

$$e_{\theta\theta} = \frac{k}{3(1 - 2\sigma)} [\tau_{\theta\theta} - \sigma(\tau_{rr} + \tau_{zz})] \quad (16)$$

$$e_{zz} = \frac{k}{3(1 - 2\sigma)} [\tau_{zz} - \sigma(\tau_{rr} + \tau_{\theta\theta})] \quad (17)$$

Equilibrium of an arbitrary volume element demands²

$$\tau_{rr} = A + \frac{B}{r^2} \quad (18)$$

and

$$\tau_{\theta\theta} = A - \frac{B}{r^2} \quad (19)$$

where A and B are constants with respect to spatial coordinates. (A

² J. Prescott, Applied Elasticity, Dover Publications, New York, 1946, p. 330.

should not be confused with the electrode area used in previous sections.) We assume

$$e_{zz} = 0 \text{ for all } (r, \theta, z) \quad (20)$$

$$e_{\theta\theta} = 0 \text{ for } r = a \quad (21)$$

and $\tau_{rr} = P \text{ for } r = b \quad (22)$

Manipulation of (15) through (22) allows the evaluation of the constants A and B as

$$A = - \frac{b^2}{a^2} \frac{P}{\frac{b^2}{a^2} + (1 - 2\sigma)}$$

and

$$B = \frac{b^2(1 - 2\sigma)P}{\frac{b^2}{a^2} + (1 - 2\sigma)}$$

Thus

$$\begin{aligned} -\frac{1}{\rho} \frac{\partial \rho}{\partial P} &= \frac{e_{rr} + e_{\theta\theta} + e_{zz}}{P} \text{ yields} \\ \frac{1}{\rho} \frac{\partial \rho}{\partial P} &= \frac{2(1 + \sigma)k}{3} \frac{b^2}{a^2} \frac{1}{\frac{b^2}{a^2} + (1 - 2\sigma)} \end{aligned} \quad (23)$$

Also

$$\frac{1}{b} \frac{\partial b}{\partial P} = \frac{e_{rr}|_{r=b}}{P} = \frac{k(1 + \sigma)}{3} \left[1 - \frac{b^2}{a^2} \right] \frac{1}{\frac{b^2}{a^2} + (1 - 2\sigma)} \quad (24)$$

Combining (2), (14), (23), and (24) we obtain

$$\frac{1}{C} \frac{\partial C}{\partial P} = \frac{2}{3} \left[\frac{(1 + \sigma) \frac{b^2}{a^2}}{\frac{b^2}{a^2} + (1 - 2\sigma)} \right] \left[\frac{(\epsilon - 1)(\epsilon + 2)}{3\epsilon} + \frac{1 - \frac{a^2}{b^2}}{2 \ln \frac{b}{a}} \right] k \quad (25)$$

E. The Case $\sigma = \frac{1}{2}$

The equations derived above reduce to the expressions one would obtain if the dielectric were considered to be a compressible fluid when σ is set equal to $\frac{1}{2}$.

Equation (8) for the parallel plate arrangement becomes

$$\frac{1}{C} \frac{\partial C}{\partial P} = \left[\frac{(\epsilon - 1)(\epsilon + 2)}{3\epsilon} + 1 \right] k \quad (26)$$

while (13) is unaltered. The difference between the two cases arises from the fact that in (13) the area was allowed to vary while in the former case it was not. The deviation of

$$\frac{1}{C} \frac{\partial C}{\partial P}$$

from (26) when $\sigma \neq 0.5$ is given by the factor

$$\frac{1}{3} \left(\frac{1 + \sigma}{1 - \sigma} \right)$$

The capacitance-pressure coefficient for the cylindrical configuration, (25), becomes

$$\frac{1}{C} \frac{\partial C}{\partial P} = \left[\frac{(\epsilon - 1)(\epsilon + 2)}{3\epsilon} + \frac{1 - \frac{a^2}{b^2}}{2 \ln \frac{b}{a}} \right] k \quad (27)$$

The deviation of

$$\frac{1}{C} \frac{\partial C}{\partial P}$$

from the value given in (27) when $\sigma \neq 0.5$ is thus given by the factor

$$\frac{2}{3} \left[\frac{(1 + \sigma) \frac{b^2}{a^2}}{\frac{b^2}{a^2} + (1 - 2\sigma)} \right]$$

III. APPARATUS

The experimental arrangement employed to investigate the validity of (8) is shown in Fig. 1.* Pressure was applied by means of a Baldwin tensile testing machine. The cell makes use of a "sandwich" arrangement wherein two disc samples (2" in diameter, 0.050" thick) of dielectric are pressed between three brass electrodes, the outer electrodes being grounded. The capacitance thus formed is well shielded and stray capaci-

* This cell was designed by C. A. Bieling.

tances are minimized. Lateral displacements are kept small by an annular ring of steatite ceramic which is in turn surrounded by a ring of Ketos steel. Pressures of 23,000 lb on the two-inch sample discs have been applied without damaging the cell.

Although measurements could be made with ease in this cell it is not without disadvantages. The steatite ring has a rather high dielectric constant which tends to increase fringing effects. These effects are furthermore pressure dependent since the electrode separation varies as pressure is applied. Also loss measurements could not be obtained as leakage along the steatite surface was larger than the leakage through the samples of the polyethylene-butyl rubber compound investigated.

An attempt was made to eliminate fringing effects by making measurements on samples of varying thickness and extrapolating to zero thickness but results were too uncertain to be of quantitative value. The uncertainty resulted from the inability to cast the sample discs with uniform thickness an effect which becomes pronounced with very thin samples. It was possible, however, to estimate the total stray capacitance in this manner and it was found to be about 10 per cent of the sample capacitance and only slightly dependent upon pressure.

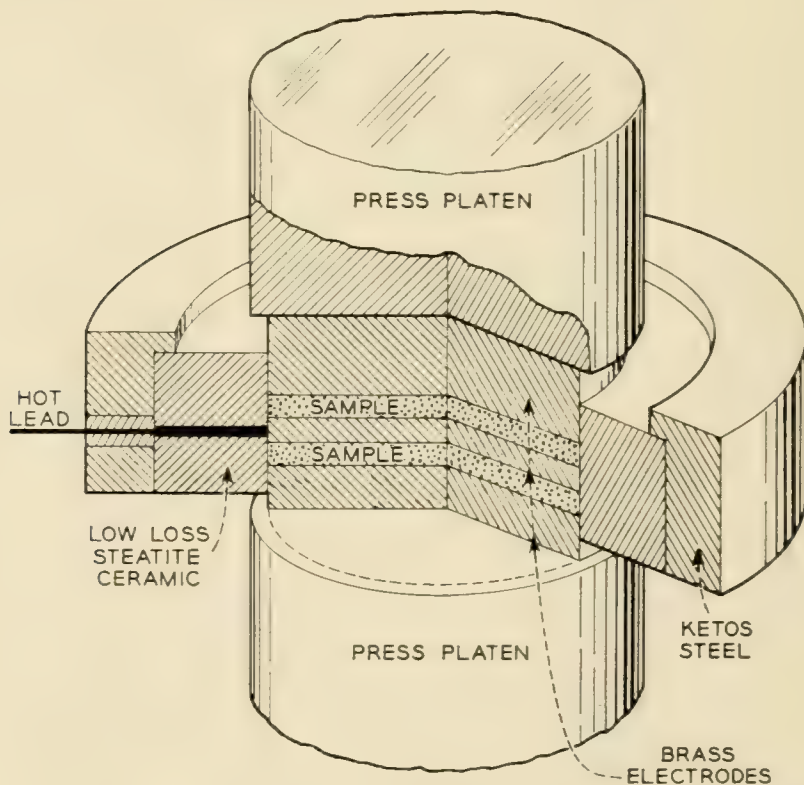


Fig. 1 — Cell employed to measure $(1/C)(\partial C/\partial P)$ with sheet specimens under one dimensional pressure.

Also, due to non-uniformity of the thickness of the specimens, it was found that the capacitance-pressure coefficient was larger at low pressures than at high pressures. This was apparently due to the initial squeezing out of voids between the electrodes and samples and the difficulty was removed by applying silver paint electrodes to the sample discs.

A Western Electric capacitance bridge was used to make capacitance measurements with this cell. The temperature was kept at 25°C and the humidity of the room was maintained at 50 per cent. The frequency used was 10 kc.

It is believed that the fact that the steatite is much more rigid than the specimens makes this experimental arrangement closely approximate the assumptions made in deriving (8) (i.e. lateral strains are negligible). It is important, however, that the specimens be cut to fit the steatite ring very closely.

Experiments which correspond to the cylindrical capacitor under biaxial stress have been performed as follows.* The specimens in this case were lengths of cable which consisted of a copper wire central conductor surrounded by the polyethylene-5 per cent butyl rubber compound. b/a was 3.87 in most of the measurements but some data were obtained for $b/a = 4.68$ (actual dimensions 0.620"/0.160" and 0.750"/0.160"). Twenty foot lengths of cable were placed in a long tank provided with a seal at one end. The end of the cable inside the tank was closed such that the center conductor was isolated. The tank was then filled with water which served as the outer conductor, tap water having sufficiently high conductivity. Pressure was applied by the water.

A Leeds and Northrup capacitance bridge was used and the measurements reported were made at 10 kc.

IV. RESULTS

It is experimentally observed in all the cases considered that plots of C versus P are nearly linear for polyethylene-5% butyl rubber. This may be shown to be in agreement with the foregoing theories as follows. Equation (8) may be written

$$\int_{C(0)}^{C(P)} d \ln C = \int_0^P \left[\frac{(\epsilon - 1)(\epsilon + 2)}{3\epsilon} + 1 \right] \frac{k}{3} \left[\frac{1 + \sigma}{1 - \sigma} \right] dP \quad (28)$$

* The investigation of radial compression on cylindrical (cable) specimens was carried out by A. W. Lebert and O. D. Grismore of Bell Telephone Laboratories.

The integrand is approximately constant so

$$\ln \frac{C(P)}{C(0)} \cong \left[\frac{(\epsilon - 1)(\epsilon + 2)}{3\epsilon} + 1 \right] \frac{k}{3} \left[\frac{1 + \sigma}{1 - \sigma} \right] P$$

and

$$C(P) = C(0) \exp \left\{ \left[\frac{(\epsilon - 1)(\epsilon + 2)}{3\epsilon} + 1 \right] \frac{k}{3} \left[\frac{1 + \sigma}{1 - \sigma} \right] P \right\}$$

$kP \cong 10^{-2}$ for the highest pressures used in the present experiments and

$$\left[\frac{(\epsilon - 1)(\epsilon + 2)}{3\epsilon} + 1 \right] \frac{1}{3} \left[\frac{1 + \sigma}{1 - \sigma} \right]$$

is of the order of unity so the exponential may be expanded as

$$C(P) \cong C(0) \left\{ 1 + \left[\frac{(\epsilon - 1)(\epsilon + 2)}{3\epsilon} + 1 \right] \frac{k}{3} \left[\frac{1 + \sigma}{1 - \sigma} \right] P \right\}$$

This treatment applies only to dielectrics for which ϵ , k , and σ are insensitive to pressure. Equations (13) and (25) may be treated similarly.

Values obtained experimentally and theoretically for polyethylene-5% butyl rubber are compared in Table I. The experimental values represent averages of many measurements. Agreement is considered adequate but more careful experiments are needed. The necessary parameters assumed in making these comparisons are:

$$\epsilon = 2.28$$

$$k = 2.14 \times 10^{-6} / \text{psi}$$

$$\sigma = 0.50$$

TABLE I — EXPERIMENTAL AND THEORETICAL VALUES FOR POLY-ETHYLENE-5 PER CENT BUTYL RUBBER

Sample	Pressure	$\frac{1}{C} \frac{\partial C}{\partial P}$ (/10 ⁶ psi)	
		Experimental	Theoretical
Sheet	one-dimensional	3.3*	3.74
Cable $\begin{cases} b/a = 3.87 \dots\dots\dots \\ b/a = 4.68 \dots\dots\dots \end{cases}$	radial	2.4	2.27
	radial	2.2	2.22

* This value has not been corrected for stray capacitance. Such a correction would tend to make the agreement between experimental and theoretical results better.

V. SUMMARY

Equations relating electrical capacitance and pressure have been derived for plane capacitors under one dimensional and hydrostatic pressures and cylindrical capacitors under radial pressure. The dielectric material has been assumed to be an elastic solid but the relationships also apply to fluid dielectrics when Poisson's ratio is set equal to $\frac{1}{2}$. Experiments corresponding to the assumptions have been described briefly and experimental results are found to be in agreement with the theoretical predictions.

The results are of practical value in making estimates of the dependence of attenuation of submarine cables on pressure. The equations may also be put in forms useful for determining the dependence of the dielectric constant on pressure from capacitance measurements.

ACKNOWLEDGEMENT

The author would like to acknowledge several valuable discussions with G. T. Kohman and A. C. Lynch. Contributions to this work were also made by A. W. Lebert and O. D. Grismore. C. A. Weatherington assisted in some of the measurements. J. A. Lewis made several important comments on the manuscript.

ERRATA

The Effects of Surface Treatments on Point Contact Transistor Characteristics by J. H. Forster and L. E. Miller, B.S.T.J., **35**, pp. 767-811, July, 1956. Figs. 3, page 776, and 10, page 787, were inadvertently interchanged.

Cable Design and Manufacture for the Transatlantic Submarine Cable System by A. W. Lebert, H. B. Fisher and M. C. Biskeborn, B.S.T.J., **36**, pp. 189-216. Table I, page 3, the material for type B armor wire should be medium steel instead of mild steel. Page 207, the equation for Z_0 should read

$$Z_0 = \frac{b}{\sqrt{\epsilon}} \log \frac{D}{b} \text{ ohms}$$

Reading Rates and the Information Rate of a Human Channel

By J. R. PIERCE and J. E. KARLIN

(Manuscript received August 31, 1956)

The limitation on the rate at which information can be transmitted over an ordinary telephone channel is a human one. In this study people read words as fast as they were able to; from these results some deductions are made about the capacity of a human being as an information channel. The discrepancy between human channel capacity measured thus (40–50 bits/sec) and telephone and television channel capacity (about 50,000 bits/sec and 50,000,000 bits/sec respectively) is provocative.

INTRODUCTION

In communication over an ordinary telephone channel, the limitation on the rate at which information can be transmitted appears to be a human one. For instance, by use of a vocoder, the required channel capacity can be reduced greatly with only a moderate reduction in the quality of the reproduced speech.¹

It would be of great interest to measure the information rate necessary to provide a satisfactory sensory input to a human being. It is not clear how this could be done. Something which may be related and for which a lower bound can be measured is the capacity of a human being as an information channel.

An evaluation of and understanding of the limitations on the information rate of the human channel might ultimately be of practical importance for two reasons. First, it might help to tell us what sort of task to set a human being when he is necessarily a part of a system involving information transmission. Thus, a man can transmit information faster by reading than by tracking. Secondly, the understanding might somewhat illuminate the problem of the channel capacity necessary to provide a satisfactory sensory input, and so might help to reduce the channel capacity required in electrical communication between human beings.

Previous investigations indicate^{2, 3} that reading aloud attains the fastest rate at which a human being can be demonstrated to transmit

information, as contrasted with, for example, typing, playing the piano, or tracking.*

The work presented here, while undertaken independently, is in general similar to and in agreement with that reported for reading rate experiments by Licklider, Stevens and Hayes,² and by Quastler and Wulff.³ However, we have considered some factors in more detail than these workers, and also, contrary to the former group, we find that, under optimal conditions, reading with tracking has a lower information rate than reading alone.

The chief problem investigated was:

- (1) Taking people as they are, with no additional training, how fast, in bits per second, can they transmit information by reading?
- (2) What principal factors control this limiting rate?

The experimental procedure consisted simply of people reading aloud as rapidly as they could typed lists of words. Each list was composed of a single vertical row of 12 groups of 5 words, giving a total of 60 words per page. In each instance, the words were chosen at random from a given vocabulary of words. If n is the number of words in the vocabulary and if the words are chosen with equal probabilities, and if all words are read correctly,† the amount of information which is conveyed or transmitted through the human being measured in bits is⁴

$$\log_2 n \text{ bits/word}$$

When the vocabulary for a particular experiment has much fewer than 60 words, certain words must necessarily be repeated several times within a list. When the vocabulary is much greater than 60 words, repetitions are necessarily few and differences in reading rate among different vocabularies would be expected only if the vocabularies differed in nature, as in syllable length or familiarity of words.

Unless otherwise specified, each result quoted below is the average reading speed for two lists for each of three readers, chosen as representing fast, medium and slow readers for people with at least a high school education. The results on these three readers are substantially similar to those on ten similar readers used in preliminary experiments. The chief experiments performed, and some interpretations of them, follow under numbered headings. Some supplementary experiments are then described briefly and the over-all results are commented on.

* Here tracking means successively pointing to a series of marks.

† In preliminary experiments the reader's voice was recorded, and it was found that errors in reading aloud occur very seldom if ever.

PRINCIPAL EXPERIMENTS

Experiment 1: Effect of Vocabulary Size

The larger the vocabulary size the higher the information rate conveyed by a given word reading rate. However, one might think that it would be possible to read randomized lists of, say, 4 words substantially faster than lists of 8, or 16 or more words.* How is the word rate affected as the vocabulary size increases?

To investigate vocabulary size as such, it is necessary as far as possible to avoid the influence of differences in word length or familiarity. To this end, words in each vocabulary were chosen at random from the 500 most common words in the language;⁵ a few words were then changed so as to keep an average of 1.5 syllables/word for each vocabulary. Figs. 1(a) and 1(b) show parts of typical lists for vocabulary sizes of 2 and 256 words respectively. The order of reading the different size vocabularies was randomized.

Fig. 2 shows that *reading rate is essentially independent of vocabulary sizes from 4 to 256 words when familiarity and word length are kept fairly constant*. The reading rates for the three readers for the 256-word vocabulary are 3.8, 3.7 and 3.0 words/sec, giving information rates of 30, 30 and 24 bits/sec respectively.

The word rate for a 2-word vocabulary is systematically a little greater than for larger vocabularies. This effect, which is statistically significant, is best seen in Fig. 2 in the average curve (dashed). The writers feel on the basis of subjective impressions that this may result from a tendency to group words in pairs in recognizing and speaking them. Among 2 words there are only 4 ordered pairs. It is apparent from the data that no such effect is noted among the 16 ordered pairs occurring with the 4-word vocabulary.

The last point on the curves in Fig. 2 illustrates the importance of familiarity and word length. When words are taken at random from a 5,000-word dictionary (12.3 bits/word), the reading rates drop to 2.8, 2.7 and 2.1 words/sec, yielding information rates of 34, 33 and 26 bits/sec respectively, which are very close to the rates 30, 30, 24 for the 256-word vocabulary.

However, these dictionary lists involve some unfamiliar words and average 2.2 syllables/word.

* When the light is very dim, the reading rate is slowed, and is faster for small vocabularies than for large vocabularies. Reading tests were done at normal light levels, which are very much brighter than those at which a slowing due to inadequate illumination is observed.

Experiment 2: Effect of Word Length and Familiarity

It was not clear from Experiment 1 how much of the drop in word rate for the dictionary list was affected by decreased familiarity and how much by increased word length.

These two variables were then untangled in a separate experiment. Word lists were prepared which kept both length and familiarity relatively constant for a given list. The words were chosen from a list of the 20,000 most frequently encountered words in the language.⁶ Reading rates were measured for the thousand most familiar words, for the ninth to tenth thousand most familiar, and for the nineteenth to twentieth thousand most familiar words.

The results are shown in Figs. 3(a), (b), and (c). There is considerable consistency among readers as to the relative effect of length and familiarity. The most familiar trisyllable words, for example, are read about as rapidly as the least familiar monosyllables.

A confirmatory demonstration of the effect of familiarity upon reading rate is shown in Fig. 4. This shows reading rates for randomized lists of eight nonsense words averaging 1.5 syllables/word (e.g., jevhin, tosp) which are necessarily totally unfamiliar when the reader first encounters them. As the reader becomes more familiar with the words on successive readings, his word rate increases until he approaches the rates of familiar words in Fig. 2.

Experiment 3: Preferred Vocabulary for Increasing Transmission Rate

The transmission rate is the product of the reading rate and the logarithm to the base 2 of the vocabulary size. To maximize the rate we

Fig. 1 — Parts of typical lists for vocabulary sizes.

grew	foot
action	tomorrow
grew	count
grew	issue
action	rain
action	month
grew	earth
grew	cook
action	build
action	corner
grew	yard
action	history
grew	forest
action	pleasant
grew	wrong
(a)	(b)

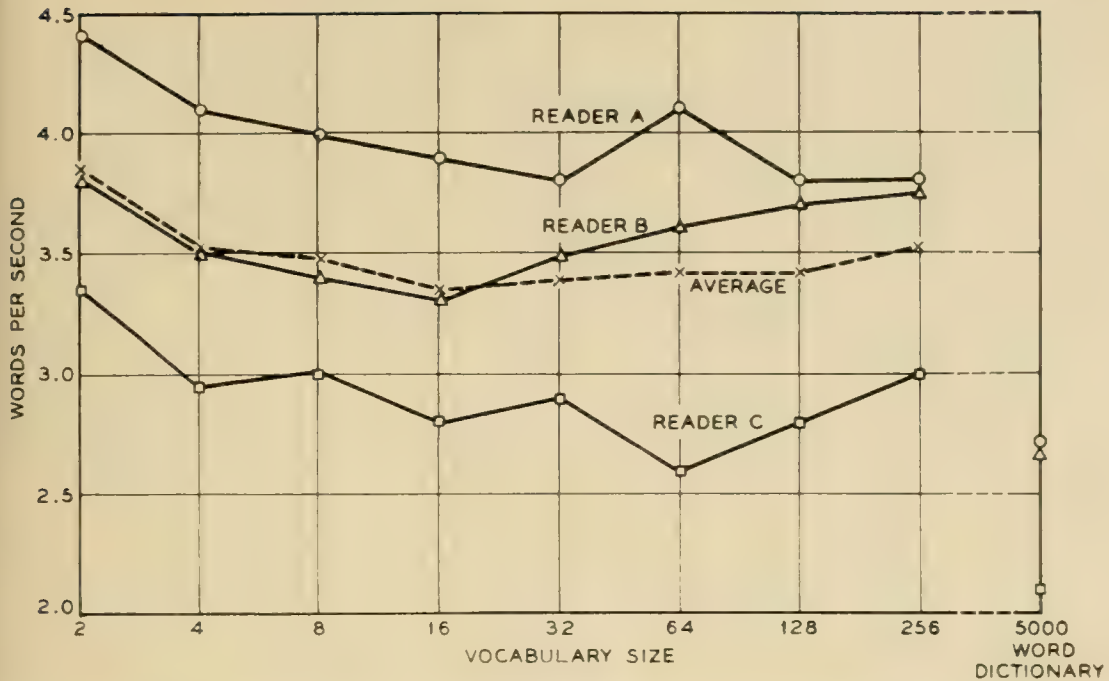


Fig. 2 — Reading rate is essentially independent of vocabulary sizes under certain conditions.

must make each of these factors large. As Fig. 3 indicates, reading speed tends to decrease as vocabulary size increases. From the data in Fig. 3, a rationale (shown in Appendix I) was developed for use in searching for an improved vocabulary which would maximize transmission rate. This indicated that the 2,500 most familiar monosyllables chosen with equal probability should form a very good vocabulary and one which is simple to construct and use. For a 2,500-word list we have 11.3 bits/word.

Reading speeds for such preferred lists were 3.7, 3.4 and 3.0 words/sec, giving information transmission rates of 42, 39 and 34 bits/sec. Some data on the distribution of this rate found among Bell Telephone Laboratories employees is given in Fig. 5.

Experiment 4: Prose and Scrambled Prose

The experiments above were all with discrete words. Reading rates for non-technical prose* are appreciably higher — 4.8, 4.7 and 3.9 words/sec for the three readers. However, such prose has a good deal of redundancy. Shannon⁷ arrives at a figure of around 1 bit/letter for a

* Extracts were taken from New York Herald Tribune, the novel "East River" by Sholem Asch, "Vermont Tradition" by Dorothy Canfield Fisher and the *Scientific American*. Such material was chosen as being of the same sort of prose as was used by Dewey⁵ in his word counts from which Shannon⁷ made his estimate of information content of printed English.

27-word alphabet including the space, or 5.5 bits/word for the average of 4.5 letters/word plus one space following a word. Newman and Gerstman⁸ give a figure of 2 bits/letter. It is quite uncertain, however, what the true value may be. Table I compares the information rate for the preferred list with that for prose assuming 5 and 10 bits/word.

When words were taken at random from the same prose sources, the reading rates dropped to 3.7, 3.3 and 2.7 words/sec. These rates are about the same as for the preferred list.

The information content of scrambled prose can be estimated much more accurately than that for prose, since the correlations associated

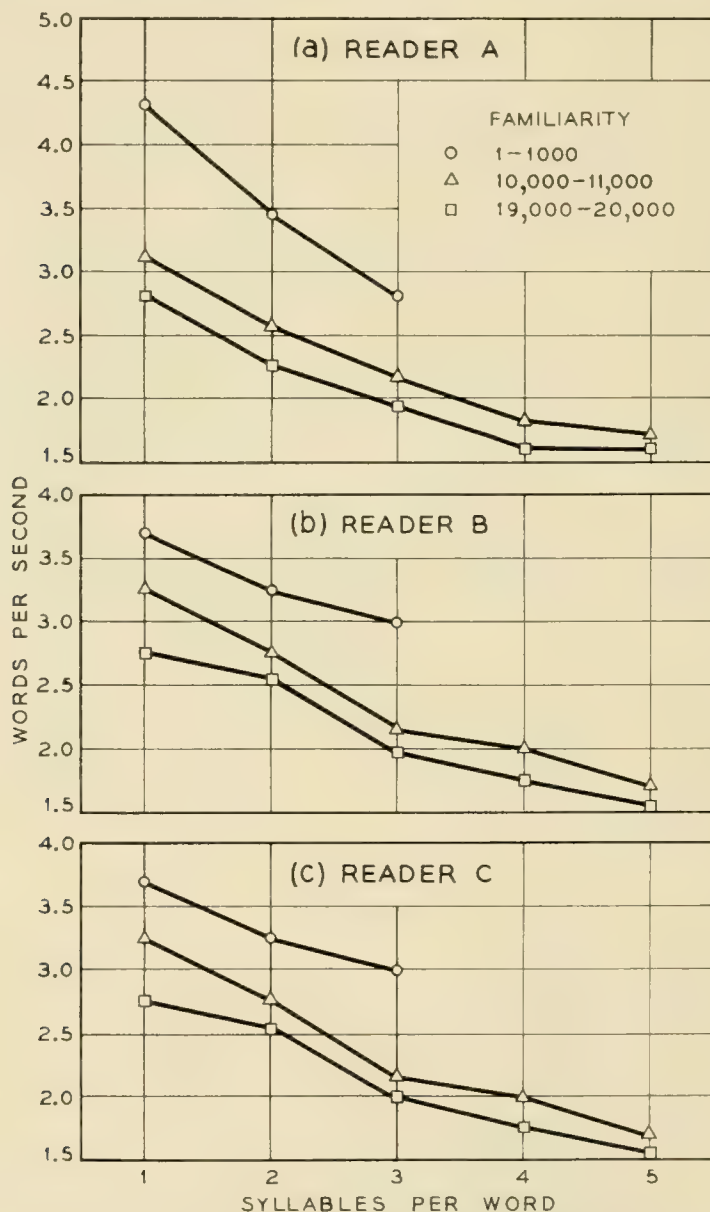


Fig. 3 — Effect of word length and familiarity.

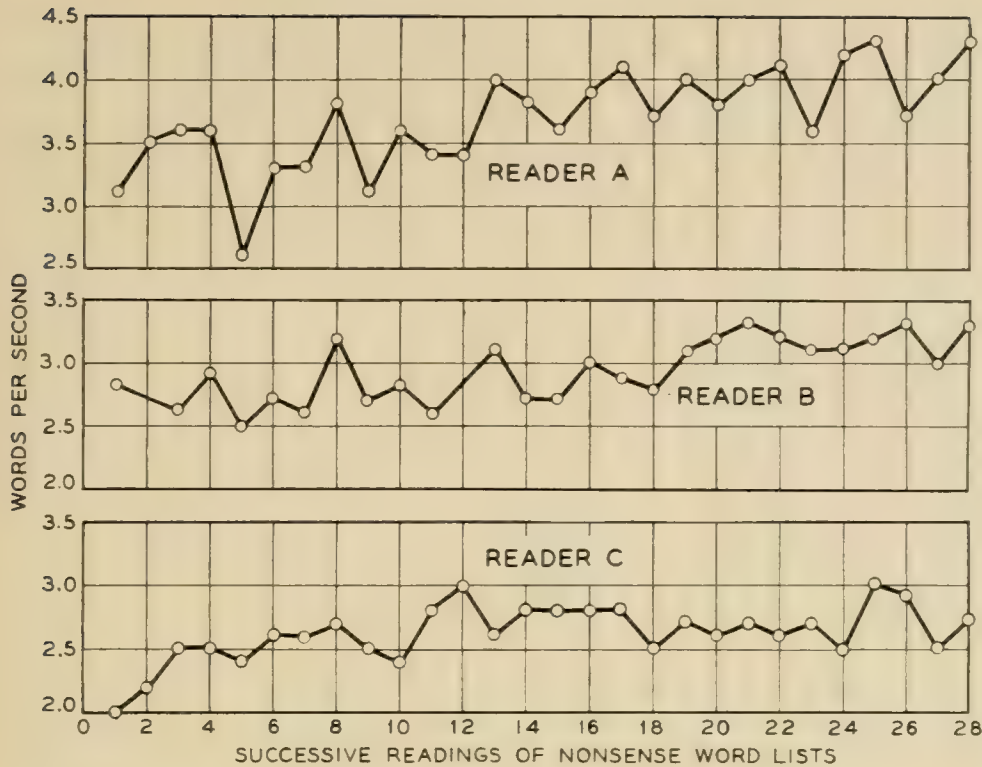


Fig. 4 — Confirmatory demonstration of the effect of familiarity upon reading rate.

with word order have been removed and the bits/word depend only on the frequency of occurrence of words in prose, which is known. Thus, Shannon⁷ gives a figure of 11.82 bits/word which applies to scrambled prose, provided the prose has the same word frequencies as that from which the statistics were derived. The information rates for words from a 5,000-word dictionary (Experiment 1) for the preferred lists, and for scrambled prose are given in Table II.

The information rate for scrambled prose is less reliable than the others, because we are not sure that the word frequencies used by Shannon apply to the prose used by us, but we used the type of material cited by the reference he quotes. It is clear that the information rate for scrambled prose is high as compared with most other lists.

Table II shows the gain which may be made by fitting the task to the human being — in this case, by choosing a suitable word list. We may note that the gain appears greater in the case of reader A than in the case of reader B. This need not be experimental error. One would suppose that there are optimal lists for individuals. Indeed, if we compare Figs. 3(a) and 3(b) we see that for reader A the word rate for monosyllables drops by a factor 0.72 in going from the first thousand to the tenth thousand, while for reader B the drop is only a factor 0.88. This

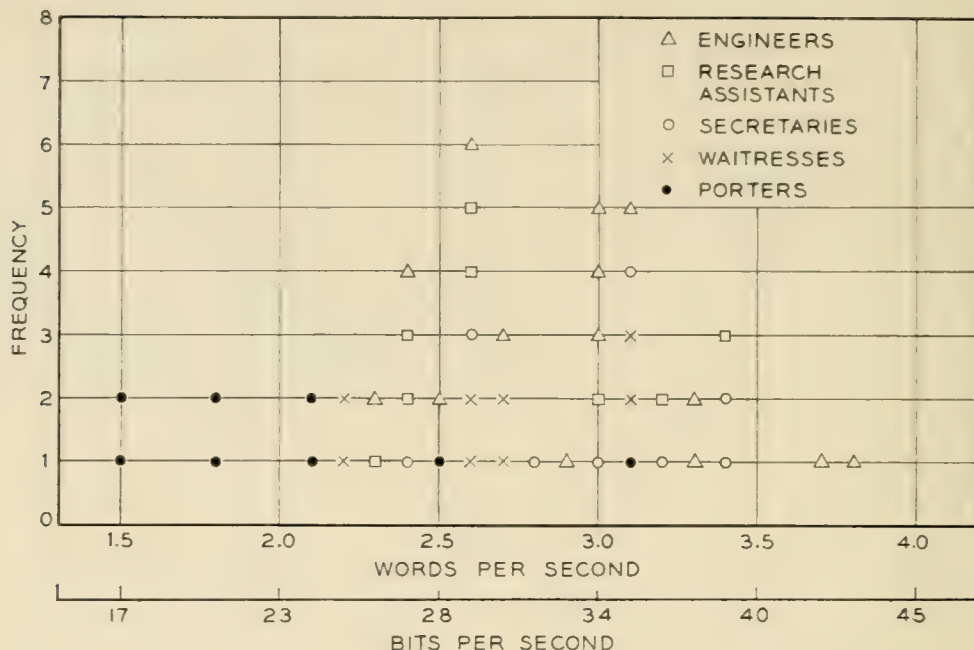


FIG. 5 — Distribution of reading rates for preferred vocabulary.

indicates that the optimal list would be somewhat different for reader A than for reader B. More extensive data would, however, be required to confirm this hypothesis.

Experiments not described here in detail showed that reading rates for digrams (successive pairs of words related as in English text) are intermediate between those for prose and discrete words.

Experiment 5: Effect of Multiple Channels

Licklider¹ has found that when the reader attempts simultaneously to perform a tracking operation while he is reading, his reading rate remains almost unimpaired, and the tracking information is added to that of reading alone. This two-channel transmission gave him his highest rate of transmission. We obtained the reverse finding. Reading the preferred list gave us our highest transmission rate. Simultaneous

TABLE I

Material	Information Rate (bits/sec)		
	A	B	C
Preferred list.....	42	39	34
Prose (5 bits/word).....	24	23	19
Prose (10 bits/word).....	48	47	39

TABLE II

Material	Information Rate (bits/sec)		
	A	B	C
5,000-word dictionary	33	33	26
Preferred list	42	39	34
Scrambled prose	43	39	32

reading and tracking gave a lower total transmission rate. However, Licklider and we agree on the magnitude of this maximum — between 40 and 45 bits/sec for facile test subjects.

Measurements on combined reading and tracking rates were made in Experiment 5 using words from the preferred lists. Whereas Licklider's readers made a dot within a box next to the word read, our readers placed a dot as close as possible to a vertical line next to the word read (e.g. dog |·). The computation of transmission rate is shown in Appendix II. The reading-while-tracking rates were 2.4, 2.0 and 1.4 words/sec. The computed information rates are given in Table III.

It may be seen that the reading rate during tracking dropped so much that the two channels together give a total information rate less than those for reading the preferred list alone. Licklider's reading lists were words chosen randomly from a dictionary and are presumably not chosen optimally for maximum information rate — his information rates for reading alone were 30–35 bits/sec, as compared with the 32–43 bits/sec found here for the scrambled prose and preferred lists. However, if we assume that our reading-while-tracking rate, which is much slower than the reading rate for scrambled prose or for the preferred lists, is limited largely by tracking, we might have obtained a slightly higher information rate in reading-while-tracking by using a larger list of words. This is suggested by the fact that Licklider's and our experiments obtain about the same reading-while-tracking speeds.

TABLE III

	Information Rate (bits/sec)		
	A	B	C
Reading (while tracking)	26.6	22.1	15.4
Tracking (while reading)	10.7	11.0	11.7
Reading and Tracking	37.3	33.1	27.1
(Rates for same word list from Experiment 3 — reading only)	(42)	(39)	(34)

Experiment 6: Effect of Physiological Utterance Limitations

One of our best indications that the maximum reading rate of a subject is determined by mental rather than by physical limitations is that discrete word lists were read no faster silently than aloud. This may appear contrary to very high silent reading rates widely quoted. This can be explained by the fact that in reading much prose we do not and need not recognize every word in order to get the sense. Presumably, if an author made every word say something, his prose could not be read with understanding at such high rates.

We can also show in another way that the mere uttering of the words does not determine the reading speeds observed. A memorized prose phrase ("This is the time for all good men to come to the aid of their country") was repeated several times at rates of 7.5, 9.1 and 8.4 words/sec for the three readers.

Fig. 6 compares word rates for repeating a phrase with the word rates previously discussed. The radically faster rate for repeating a phrase is not the only feature to be observed in this figure; the three readers are not in the same order of speed as is preserved through the reading experiments. This would suggest that it is word recognition rather than speaking speed which accounts for differences among the reading rates of different people.

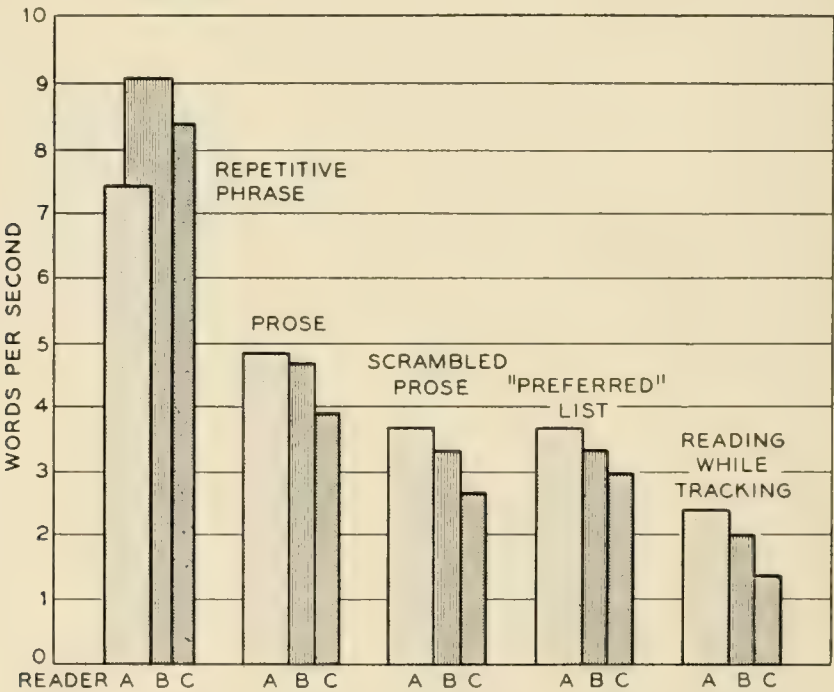


Fig. 6 — Effect of physiological utterance limitations.

DISCUSSION AND SUPPLEMENTARY EXPERIMENTS

Conclusions from Principal Experiments

The conclusions which can be reached with reasonable assurance from these experiments are rather narrow. They might be stated:

1. Information is best transmitted through a human channel by means of well-chosen acts (reading well chosen words in this case) involving many bits per act, that is, much choice per act. Cutting down drastically the bits per act does not substantially increase the speed at which the individual act is accomplished.

2. The lower bound of information transmission through the human channel of rapid readers seems to be about 43 bits/sec. This estimate is a little higher than that found by Licklider,² and may be close to a limiting rate.

3. This limiting rate can be achieved by the simple act of reading either randomized lists from suitably selected words or scrambled prose.

4. Both familiarity and length of words are important in determining reading speed. The relative effect of these two variables on reading speed is rather complex.

Beyond these narrow conclusions, there is much understanding yet to be achieved in the general field of the speed of human mental and physical responses and operations. Thus, it seems worth while to mention other experiments which were done in the course of the present investigation and experiments carried out by other workers, and to speculate somewhat concerning the whole of this experimental work.

Multiple Tasks

The reading-while-tracking experiments touch on an important problem. We have all heard of wireless operators who can receive and subsequently type out a message while carrying on a conversation or playing chess. There is nothing in this feat to indicate an information rate greater than that we have found. Actually, the rate of receiving prose by International Morse Code by ear is around 0.58 word/sec;⁹ this is slow compared with the rates we have considered.

Our experiments with tracking followed experiments in which words in the lists were randomly printed in red or black, and in which the subject spoke red words in a louder tone of voice than black words, or pressed one key for red words and another for black words. In these cases, the added information, one bit per word, was so small as to make no clearly discernible difference in information rate for the large vo-

cabularies. The speed for reading loud and soft was less than for reading-while-keying. This may imply something about the relative efficiency of human beings performing two tasks by using two sets of muscles as against using one set in two different ways.

It is common experience that we can walk about and carry out other simple tasks while talking or thinking. It is possible though not obvious that some sort of automatic, almost purely reflexive response — as, moving the left hand when the right hand is touched — could with practice be carried out quite independently of a task such as reading. The information rate for such responses would be small, the experimental error would make it difficult to settle the question, and the interpretation of such an experiment would not be entirely clear.

The Patterns Which Govern Reading Time

Early in the experiments the question was raised whether readers may not read letter by letter or syllable by syllable. Several findings bear on this.

Fig. 3 shows clearly that the reading time for a two-syllable word is much less than twice the reading time for a one-syllable word.

One of us knows a negligible amount of German. German syllables are, however, reasonably familiar. It was found that in reading German aloud he had the same reading rate in syllables per second as a man whose native language was German had in words per second. The two readers had substantially the same reading speed in English. Presumably in reading German one man recognized syllables and the other recognized words. This also reinforces the conclusion that reading rate is not limited by the time taken to utter words.

Some experiments were done using lists of common Chinese characters and lists of the corresponding English words. Average word rates over three lists for two readers who could read both languages are given in Table IV. The slightly lower rate for English is plausibly explained by the fact that Chinese was the reader's native language. All words were

TABLE IV

Reader	Words/sec	
	Chinese Words	English Words
E	2.7	2.3
F	3.3	3.2

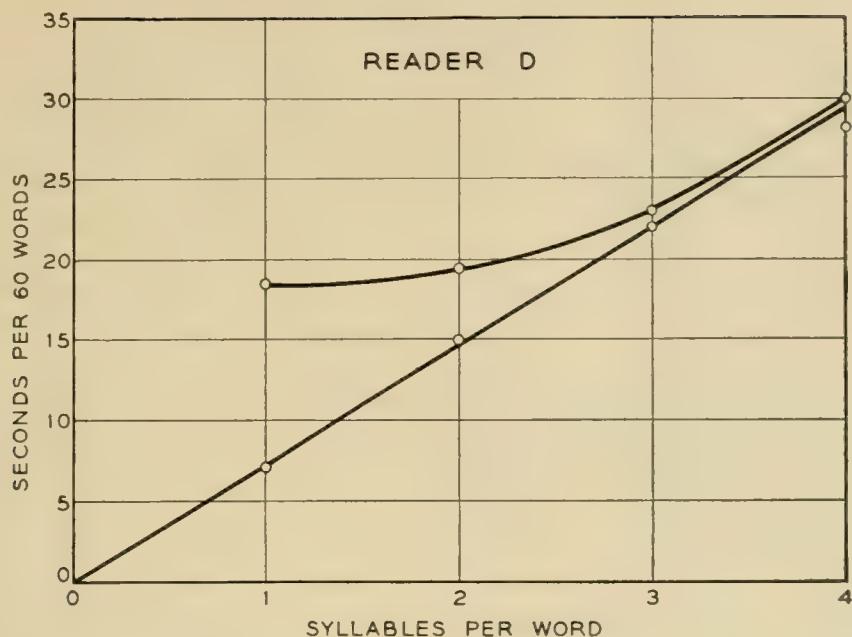


Fig. 7 — Patterns governing reading time.

necessarily monosyllables in Chinese and happened to be monosyllables in English.

In one case a word is made up of a sequence of letters, each standing for a sound, and in the other it is made up of a number of strokes which are meaningless individually, yet in each case a word is taken as a unit or pattern requiring nearly the same time for reading.

We found that the rate for reading arabic numerals is substantially the same as for reading familiar words. Each numeral is an individual pattern to be recognized. •

In a first effort to find the effect of syllable length on reading rate, a subject read several lists made up respectively from vocabularies of 16 single-syllable, 16 two-syllable, 16 three-syllable, and 16 four-syllable words. None of the words was very unfamiliar to start with, and all were presumably very familiar after the subject had read several randomized lists composed of the same words.

The outcome of the experiment is shown in Fig. 7. For comparison, the points associated with the lower straight line are *time for 60 words* for repeating, as rapidly as possible, a one-, a two-, a three- and a four-syllable word. The points on the upper curve are reading time for 60 words for the randomized lists of the familiar one-, two-, three- and four-syllable words.

In dealing with such groups of highly and uniformly familiar words, it appears that, roughly, a certain time is required to recognize the word regardless of length, and this time governs the reading rate up to the

TABLE V

Reader	Words/sec		
	Scrambled Prose	Scrambled Paragraph	Prose
A	3.7	4.0	4.8
B	3.3	3.7	4.7
C	2.7	2.9	3.9

point at which the reader is uttering words continuously as fast as he can. This is consistent with a strong subjective feeling that what limits the rate is the difficulty of "recognizing" the word as one looks at it, and that once the word is recognized one can utter it while recognizing the next word.

It would of course be wrong to conclude from this experiment that multisyllable words are in general recognized as quickly as single syllable words, for it would be possible to recognize one among a known group of 16 multisyllable words without looking at the whole word. Indeed, Fig. 3 indicates a substantial difference of reading rate between one- and two-syllable words of like frequency of occurrence. This was not observed in reading the specially familiar lists of one- and two-syllable words.

Why is prose read faster than scrambled prose? It might be that some short phrases are recognized as individual patterns. However, there is another factor at work. A scrambled paragraph of prose is read slower than the same paragraph in its natural word order but faster than scrambled prose from a book or a long stretch of prose, as can be seen from Table V.

It should be noted that reading speed differs for different prose, and that when comparisons among prose, scrambled paragraphs and scrambled prose are made, similar material should be used.

The fact that a scrambled paragraph is read faster than scrambled prose might be explained by saying that we expect, we are more ready to recognize, words which are repetitions of earlier words or words which are closely related in sense to earlier words than we are unrelated words. Thus, the greater reading speed for prose than for scrambled prose seems to be due only in part if at all to the recognition of phrases rather than words as individual patterns.

Rate of Mental Processes

The rate at which information passes through a human channel in reading experiments is indisputable. Quastler³ has attempted to go

beyond this and estimate information processing rates in the brain from the performance of lightning calculators, by dividing the performance of the calculation into a sequence of tasks equivalent to consulting memorized multiplication tables and performing additions. It is hard to interpret such a study clearly, for it is quite possible that there are many sorts of mental acts which take different times to perform, just as multiplication and addition take different times in an electronic computer. A tentative experiment we performed indicated something of the sort.

Randomized lists were made up from vocabularies (a) of names of common animals and vegetables in equal numbers, and (b) of common men's and women's names in equal numbers. In reading these, a subject was asked, not to read the word aloud, but merely to press one key with his right and another key with his left hand; in (a) left-animal, right-vegetable; in (b) left-man, right-woman. The same subject later read the lists aloud. Pressing keys took 40 per cent longer than reading aloud. (The additional time is not related to the keying operation itself; for a 2 word list, for example, keying speed is much faster than reading speed.) Presumably an additional mental operation was involved, but it was not one for which the time was equal to that for reading. This experiment was not pursued further, partly because no clear conclusion could be drawn from it. Had it been pursued and randomized lists of the same words used repeatedly, the rate might have gone up. Conceivably, cow and horse could become for a subject merely different ways of spelling left, and lettuce and carrot variant spellings of right. In this case we would end with a two-word reading experiment.

Reading Rate as a Psychometric Datum

It is interesting to speculate on the possible relationship of reading rate to general intelligence or some other aptitude. Certainly, Fig. 5 indicates some such relationship. We might also ask in connection with Fig. 3, does a rate which falls less rapidly with frequency of occurrence indicate a larger vocabulary, and can we measure vocabulary by reading speed tests? Certainly, measuring the speed of reading aloud is very simple, and such tests might have some psychometric utility.

The Channel Capacity Required for Satisfactory Communication

In conclusion, we cannot help but wonder that the highest information rate noted — 43 bits/second — is so much lower than the channel capacity* of a telephone or a television circuit (around 50 thousand

* This is the limiting channel capacity given by (2.1) of Appendix II. The practical rate at which binary digits can be sent over a telephone circuit with simple equipment is less than 1000 bits/second.

bits/second for telephone and 50 million bits/second for TV). This would not be surprising if the limitation we observe had been one of the speeds at which words can be uttered, but it appears rather to be a mental one, one of recognizing what is before the eyes. To the authors, it seems reasonable that this mental limitation may apply to a human being's ability to absorb information, that is, to the information rate needed to present a satisfactory sensory input to a human being. If it does, then why do we need so much channel capacity to convey to him an acceptable sound or picture?

This can be explained in part by the inefficiency of our present communication methods. Despite its present imperfections, the vocoder makes it clear that clearly understandable speech can be transmitted using far less channel capacity than that required in ordinary telephony.¹

However, it is quite likely that even with the most efficient of encoding means we will have to use far more than 43 bits/second for a picture transmission channel. While only a portion of the image of the transmitted picture falls on the fovea at any instant, we can cast our eyes on any portion of the received picture. If the pick-up camera device and the received picture followed eye movements, a much less detailed picture would serve. Even with our eyes fixed, we can concentrate our attention on a particular part of our field of vision, and this is something that the pick-up camera cannot track. There may be similar effects in our apprehension of sounds.

In the light of present knowledge it is impossible to estimate the minimum channel capacity required to transmit sound and pictures in a satisfactory manner. It will take work far beyond the measurement of reading rates to enable us to make such an estimate.

ACKNOWLEDGMENTS

The writers are indebted to D. L. Letham, who carried out some work preliminary to that reported here, to J. L. Kelly for helpful comments, to D. Slepian for help in putting the appendices in a mathematically more acceptable form, to Miss Renee Hipkins who carried out most of the experiments, and to three indefatigable readers, Mrs. Mary Lutz, S. E. Michaels and A. P. Winnicky.

APPENDIX I

ON OBTAINING GOOD VOCABULARIES

The experiments in the body of the paper indicate that the choice of a good vocabulary is important in attaining a high information rate in reading lists of words.

In these experiments the randomized lists may be regarded as an information source consisting of a sequence of code elements or symbols (words) in which there is no correlation between successive symbols. If in making up the lists the s th word of the vocabulary is used with a normalized probability p_s , the entropy in H in bits per word, and hence the amount of information per word, is

$$H = - \sum_s p_s \log_2 p_s \text{ bits} \quad (1.1)$$

If all words appear with an equal probability $1/m$ where m is the number of words in the vocabulary, as in the case of experiments 1-3, p_s is $1/m$ for each of the m words and in this special case

$$H = \log_2 m \quad (1.2)$$

In the case of scrambled prose, for instance, the probabilities are different for different words. This will be true also if in making up word lists we choose words randomly from a box containing different numbers of different words.

Let t_s be the time taken to read the s th word of the vocabulary. *Let us assume that t_s is the same for the s th word no matter what context that word appears in in the randomized list.* If this is so, the average reading time per word, $\bar{t}$, will be

$$\bar{t} = \sum_s p_s t_s \quad (1.3)$$

the word rate will be $1/\bar{t}$, and the information rate R will be

$$R = \frac{H}{\bar{t}} = \frac{\sum_s p_s \log_2 p_s}{\sum_s p_s t_s} \quad (1.4)$$

Suppose we have available a vocabulary of words and know the reading time t_s for each word. The problem is to choose p_s in terms of t_s as to maximize R . This is easily done; however the result can also be obtained as a special case of the problem treated in Appendix 4 of "The Mathematical Theory of Communication."⁴ In Shannon's $\ell_{ij}^{(s)}$, the subscripts i, j refer to passing from state i to state j . In our case there is only one state, and $\ell_{ij}^{(s)}$ should be identified with t_s for all i and j . Similarly, we identify $p_{ij}^{(s)}$ with p_s . C is the maximum rate, so $\log_2 W = C$. Shannon's equation

$$p_{ij}^{(s)} = \frac{B_j}{B_i} W^{-\ell_{ij}^{(s)}}$$

becomes

$$p_s = 2^{-Rt_s} \quad (1.5)$$

since there is only one B .

Shannon's determinantal equation

$$\left| \sum_s W^{-t_{ij}^{(s)}} - S_{ij} \right| = 0$$

becomes

$$\sum 2^{-Rt_s} = 1 \quad (1.6)$$

In (1.5) and (1.6) we have a means of evaluating p_s in order to attain the maximum information rate R .

The data we actually have concerning words is that for some class s of words, say, the monosyllables in the 8,000-9,000 words in order of familiarity, the reading time has some value t_s , presumed to be the same for all words in the class, and that there are N_s words in this class. In this case we must assign to each word in the s th class the same probability p_s given by

$$p_s = 2^{-Rt_s} \quad (1.7)$$

and we must have

$$\sum_s N_s p_s = \sum_s N_s 2^{-Rt_s} = 1 \quad (1.8)$$

Using the same amount of data given in Fig. 3, for the 20,000 most common words, but for a different reader, estimates were made of N_s and t_s for all the classes consisting of words of each number of syllables in each range of occurrence of 1,000 words. Then the optimum values of p_s for words in each class and the maximum rate R were computed.

Using (1.4), rates were also computed for choosing words with equal probability from among the first m thousand words and from among the first m thousand monosyllables, as functions of m . These rates had

TABLE VI

Nature	Computed Rate, bits/sec
Maximum Rate.....	33.5
Maximum for equi-probability monosyllables (from first 8,000 words).....	32.4
Maximum for equi-probability among words of all lengths (first 5,000 words).....	30.2

maxima for vocabularies of optimal sizes. Table VI compares the various rates computed.

As it is much easier to make up lists from the 2,500 monosyllables among the first 8,000 words with equal probabilities than it is to make up lists from among all words with a different probability for each class, and as the information rates computed were close together, the former alternative was chosen.

The use of scrambled prose provided an easy way to make up good lists.

APPENDIX II

TRACKING EXPERIMENT

A well-known formula for channel capacity R in bits/sec is⁴

$$R = B \log_2 \left(1 + \frac{P_s}{P_n} \right) \quad (2.1)$$

This gives the limiting rate at which information can be transmitted over a channel with a bandwidth B by a signal of power P_s , in the presence of a gaussian noise of power P_n , with an error rate smaller than any assignable number.

In most cases, the actual rate is much smaller than this limiting rate. In general, the rate is the entropy of the received signal minus the entropy of the noise. In the particular case of a gaussian signal source as well as a gaussian noise, each represented by $2B$ samples a second, the calculation based on entropies gives exactly (2.1). Let us then apply (2.1) to the tracking experiment.

Suppose that a large number N of samples do have a gaussian distribution of mean square amplitude $\overline{x^2}$. Suppose that we make an error d_n in reproducing the n th sample, that these errors are gaussian, and that the mean square error is $\overline{d^2}$

$$\overline{d^2} = \frac{1}{N} \sum_n \overline{d_n^2}$$

We see from (2.1) that ideally we can use these reproduced samples to transmit M bits of information where

$$M = \frac{N}{2} \log_2 \left(1 + \frac{\overline{x^2}}{\overline{d^2}} \right) \quad (2.2)$$

In the reading and tracking experiments, randomized words from the 2,500 commonest monosyllables were arranged with equal vertical

spacings but with various horizontal positions. To the right of each word was a short vertical line. The distances x_n of these lines from the vertical centerline of the paper were obtained from a list of random numbers with a gaussian distribution such that for the list $\overline{x^2} = 1$ inch. Of course, $\overline{x^2}$ for each list would depart from this value. As the words were read, the reader used a pencil to make a dot as near as possible to the corresponding vertical line. For each sheet, the departures d_n from the vertical lines in inches were measured and $\overline{d^2}$ was computed. The number of bits M for pointing for that sheet were then taken as

$$M = \frac{N}{2} \log_2 \left(1 + \frac{1}{\overline{d^2}} \right) \quad (2.3)$$

REFERENCES

1. E. E. David, Naturalness and Distortion in Speech Processing Devices, Jour. Acoustical Soc. Am., July, 1956.
2. J. C. R. Licklider, K. N. Stevens, J. R. M. Hayes, Studies in Speech, Hearing and Communication, Technical Report, Acoustics Laboratory, MIT, Sept. 30, 1954.
3. H. Quastler et al, Human Performance in Information Transmission, Report No. R-62, Control Systems Laboratory, University of Illinois, March, 1955.
4. C. E. Shannon and W. Weaver, The Mathematical Theory of Communication, University of Illinois Press, 1949.
5. G. Dewey, Relative Frequency of English Speech Sounds, Harvard University Press, 1923.
6. E. L. Thorndike, A Teacher's Word Book of the Twenty Thousand Words Found Most Frequently, Teachers College, Columbia University, 1932.
7. C. E. Shannon, Prediction and Entropy of Printed English, B.S.T.J., **30**, pp. 50-64, Jan., 1951.
8. E. B. Newman and C. J. Gerstman, A New Method for Analyzing Printed English, J. Exp. Psychol., **44**, pp. 114-125, 1952.
9. Keith Henney, *Radio Engineers Handbook*, 3rd Ed., p. 568, McGraw-Hill, 1941.

Binary Block Coding

By S. P. LLOYD

(Manuscript received March 16, 1956)

From the work of Shannon one knows that it is possible to signal over an error-making binary channel with arbitrarily small probability of error in the delivered information. The effects of errors produced in the channel are to be eliminated, according to Shannon, by using an error correcting code. Shannon's proof that such codes exist does not provide a practical scheme for constructing them, however, and the explicit construction and study of such codes is of considerable interest.

Particularly simple codes in concept are the ones called here close packed strictly e -error-correcting (the terminology is explained later). It is shown that for such a code to exist, not only must a condition due to Hamming be satisfied, but also another condition. The main result may be put as follows: a close-packed strictly e -error-correcting code on n , $n > e$, places cannot exist unless e of the coefficient vanish in $(1 + x)^e(1 - x)^{n-1-e}$ when this is expanded as a polynomial in x .

I. INTRODUCTION

In this paper we investigate a certain problem in combinatorial analysis which arises in the theory of error correcting coding. A development of coding theory is to be found in the papers of Hamming¹ and Shannon²; this section is intended primarily as a presentation of the terminology used in subsequent sections.

We take $(0, 1)$ as the range of binary variables. By an n -word we mean a sequence of n symbols, each of which is 0 or 1. We call the individual symbols of an n -word the *letters* of the n -word. We denote by B_n the set consisting of all the 2^n possible distinct n -words. The set B_n may be mapped onto the vertices of an n -dimensional cube, in the usual way, by regarding an n -word as an n -dimensional Cartesian coordinate expression. The *distance* $d(u, v)$ between n -words u and v is defined to be the number of places in which the letters of u and v differ; on the n -cube, this is seen to be the smallest number of edges in paths along edges between the vertices corresponding to u and v . The *weight* of an n -word u

is the number of 1's in the sequence u ; it is the distance between u and the n -word $00 \cdots 0$, all of whose letters are 0.*

A *binary block code of size K on n places* is a class of K nonempty disjoint subsets of B_n where in each of the K sets a single n -word is chosen as the *code word* of the set.† Each such set is the *detection region* of the code word it contains, and we shall say that any n -word which falls in a detection region *belongs to* the code word of the detection region. The set consisting of those n -words which do not lie in any detection region we call *limbo*.‡ A *close packed code* is one for which limbo is an empty set; i.e., a code in which the detection regions constitute a partition (disjoint covering) of B_n .

A *sphere of radius r centered at n -word u* is the set $\{v: d(u, v) \leq r\}$ of n -words v which differ from u in r or fewer places. A binary block code is *e -error-correcting* if each detection region includes the sphere of radius e centered at the code word of the detection region. We say that a binary block code is *strictly e -error-correcting* if each detection region is exactly the sphere of radius e centered at the code word of the detection region.

This paper is devoted to the consideration of close packed strictly e -error-correcting binary block codes. We shall refer to such a code as an *e -code*, for brevity. Hamming¹ observes that a necessary condition for the existence of an e -code on n places is that

$$1 + n + \frac{1}{2}n(n-1) + \cdots + \binom{n}{e} \quad (1)$$

be a divisor of 2^n . In this paper we derive an additional necessary condition. Our condition includes as a special case a condition of Golay⁴ for the existence of e -codes of group type, and applies to all e -codes, whether or not they are equivalent to group codes.§

* If B_n is regarded as a subset of the real linear vector space consisting of all sequences $\alpha = (\alpha_1, \alpha_2, \cdots, \alpha_n)$ of n real numbers, then the "weight" of an n -word is simply the ℓ_1 norm (defined as $\|\alpha\|_1 = \sum_1^n |\alpha_n|$), and our "distance" is the metric derived from this norm.

† The term "block code", due to P. Elias, serves to distinguish the codes of fixed length considered here from the codes of unbounded delay introduced by Elias, Reference 3.

‡ In a communications system² using such a code, the transmitter sends only code words. If, due to errors in handling binary symbols, the receiver delivers itself of an n -word other than a code word then: (a) if the n -word lies in a detection region, one assumes that the code word of the detection region was intended; (b) if the n -word lies in limbo, one makes a note to the effect that errors have occurred in handling the word but that one is not attempting to guess what they were.

§ The terms "group alphabet" (Slepian⁵), "systematic code" (Hamming¹), "symbol code" (Golay⁴), "check symbol code" (Elias³), "parity check code", are roughly synonymous. More precisely, a group code is a parity check code in which all of the parity check forms are homogeneous ("even"), so that $00 \cdots 0$ is one of the code words; see Reference 5.

II. DISTRIBUTION OF CODE WORDS

Suppose an e -code on n places is given. Let us inquire as to the distribution of weights of code words. We denote by ν_s the number of code words of weight s , $0 \leq s \leq n$, and by

$$G(x) = \sum_{s=0}^n \nu_s x^s \quad (2)$$

the generating function for these numbers, with x a complex but otherwise free variable. We show in this section that $G(x)$ satisfies a certain inhomogeneous linear differential equation of order e .

If there exists an e -code on n places then this differential equation will have $G(x)$ as a *polynomial* solution; the necessary condition for the existence of an e -code on n places given in Section 4 is essentially a restatement of this fact.* First, however, we must derive the differential equation and obtain its solutions.

If w_α is a code word of the given e -code ($1 \leq \alpha \leq K$), define the set of j -neighbors of w_α as the set of n -words which lie at distance exactly j from w_α ; designate this set by $S_j(w_\alpha)$. ($S_0(w_\alpha)$ is the set whose only element is w_α itself.) Our derivation is based on the observation that, in an e -code on n places,

$$\bigcup_{\alpha=1}^K \bigcup_{j=0}^e S_j(w_\alpha) = B_n \quad \dagger \quad (3)$$

is a partition of B_n . For, the detection regions:

$$\bigcup_{j=0}^e S_j(w_\alpha), \quad 1 \leq \alpha \leq K$$

are disjoint, and in each such sum representing a detection region the summands are disjoint (the distance function being single valued). Furthermore, each n -word of B_n lies in some detection region (close packed property) and hence appears in one of the sets $S_j(w_\alpha)$ for some α and for some j satisfying $0 \leq j \leq e$.

The set

$$\bigcup_{\alpha=1}^K S_j(w_\alpha)$$

* The author is not yet able to demonstrate the converse. That is, suppose one obtains a polynomial solution $G(x)$ of (11), below, satisfying appropriate boundary conditions, and from it some coefficients ν_s , $0 \leq s \leq n$. It does not follow from the methods of this article that there is actually some e -code on n places for which these ν_s represent the number of code words of weight s .

† $\bigcup$ = set union.

consists of the n -words which are j -neighbors of some (not specified) code word; let us refer to these n -words simply as j -neighbors. Denote by $\nu_{j,s}$ the number of j -neighbors which are of weight s (with $\nu_{0,s} = \nu_s$, as above). Applying (3) to the n -words of weight s , we see that

$$\nu_s + \nu_{1,s} + \cdots + \nu_{e,s} = \binom{n}{s}, \quad 0 \leq s \leq n \quad (4)$$

is the total number of n -words of weight s . If we multiply (4) by x^s and sum on s , we have

$$G(x) + G_1(x) + \cdots + G_e(x) = (1 + x)^n \quad (5)$$

where

$$G_j(x) = \sum_{s=0}^n \nu_{j,s} x^s \quad (6)$$

is the generating function (with respect to s) for the numbers $\nu_{j,s}$.

We now express $G_j(x)$, $0 \leq j \leq e$, in terms of $G(x)$. Suppose code word w is of weight s ; that is, w consists of s ones and $n - s$ zeros in some order. A j -neighbor of w is obtained by choosing j places out of n and changing the letters of w in these places, 0's to 1's and 1's to 0's. If, in this procedure, q of the 1's of w are changed to 0's, so that $j - q$ of the 0's are changed to 1's, then the resulting j -neighbor of w is of weight $s - q + (j - q)$. Now, there are $\binom{s}{q}$ ways of choosing q places among the s where the letters of w are 1, and there are independently $\binom{n-s}{j-q}$ ways of choosing $j - q$ places among the $n - s$ where the letters of w are 0. Thus, of the $\binom{n}{j}$ different j -neighbors of w , the number $\binom{s}{q} \binom{n-s}{j-q}$ are of weight $s + j - 2q$. We may regard each of these as contributing $1 \cdot x^{s+j-2q}$ to the generating function $G_j(x)$ of (6) (provided $0 \leq j \leq e$, so that there is no overlap); hence, summing over all j -neighbors of a code word and then over all code words,

$$G_j(x) = \sum_{s=0}^n \nu_s \sum_{q=0}^{\infty} \binom{s}{q} \binom{n-s}{j-q} x^{s+j-2q} \quad 0 \leq j \leq e^* \quad (7)$$

From the easily verified polynomial identity

$$(x + y)^s (1 + xy)^{n-s} = \sum_{j=0}^{\infty} y^j \sum_{q=0}^{\infty} \binom{s}{q} \binom{n-s}{j-q} x^{s+j-2q}$$

* The limits $(0, \infty)$ on the q summation are merely for convenience; the binomial coefficients vanish outside the proper range, under the usual convention.

(n, s integers, $0 \leq s \leq n$) it follows that

$$\sum_{q=0}^{\infty} \binom{s}{q} \binom{n-s}{j-q} x^{s+j-2q} = \frac{1}{2\pi i} \int_C \frac{(x+y)^s (1+xy)^{n-s}}{y^{j+1}} dy$$

where contour C is, say, a small circle around the origin, taken positively. Thus

$$\begin{aligned} G_j(x) &= \sum_{s=0}^n \frac{\nu_s}{2\pi i} \int_C \frac{(x+y)^s (1+xy)^{n-s}}{y^{j+1}} dy \\ &= \frac{1}{2\pi i} \int_C \frac{(1+xy)^n}{y^{j+1}} G\left(\frac{x+y}{1+xy}\right) dy \\ &\equiv L_j G(x) \end{aligned} \quad (8)$$

where the operator L_j is thus defined. Change of integration variable gives

$$\begin{aligned} L_j G(x) &= \frac{(1-x^2)^{n+1}}{2\pi i} \int_{C_x} \frac{G(z) dz}{(1-xz)^{n-j+1} (z-x)^{j+1}} \\ &= \frac{(1-x^2)^{n+1}}{j!} \frac{\partial^j}{\partial z^j} \frac{G(z)}{(1-xz)^{n-j+1}} \Big|_{z=x} \\ &= \sum_{p=0}^j \binom{n-p}{j-p} \frac{x^{j-p} (1-x^2)^p}{p!} \frac{d^p G(x)}{dx^p} \end{aligned} \quad (9)$$

(with C_x a small circle enclosing x but not x^{-1} , $x^2 \neq 1$). Thus L_j may be regarded as a linear differential operator of order j , ($L_0 \equiv 1$).

Using this result, (5) may be given the form

$$\begin{aligned} (1+x)^n &= [L_0 + L_1 + \dots + L_e] G(x) \\ &= \frac{1}{2\pi i} \int_C \frac{y^{e-1} - 1}{1-y} (1+xy)^n G\left(\frac{x+y}{1+xy}\right) dy \\ &\equiv MG(x) \end{aligned} \quad (10)$$

this last expression as a definition of operator M . Written as a differential equation, (10) is

$$\sum_{p=0}^e \frac{(1-x^2)^p}{p!} \sum_{r=0}^{e-p} \binom{n-p}{r} x^r \frac{d^p G(x)}{dx^p} = (1+x)^n \quad (11)$$

It is straightforward that the only singularities of this equation are regular singularities⁶ at $x = \pm 1, \infty$.

III. THE DIFFERENTIAL EQUATION

In this section we discuss (11) without reference to the fact that $G(x)$ is supposed to be a generating function. That is to say, with n and e

fixed but arbitrary non-negative integers, we denote by

$$G(x) = \sum_{s=0}^{\infty} \nu_s x^s \quad (12)$$

any solution of (11) regular in the unit circle.

It proves convenient to introduce certain functions $f_{n,\xi}(x)$ defined by

$$\begin{aligned} f_{n,\xi}(x) &\equiv (1+x)^\xi (1-x)^{n-\xi} \\ &= \sum_{s=0}^{\infty} \varphi_s(n, \xi) x^s \end{aligned} \quad (13)$$

where the coefficients $\varphi_s(n, \xi)$ are given by

$$\varphi_s(n, \xi) = \sum_{r=0}^s (-1)^r \binom{n-\xi}{r} \binom{\xi}{s-r} \quad (14)$$

Here, ξ is to be regarded as a free complex variable. By $(1+x)^\xi (1-x)^{n-\xi}$ we mean $\exp(\xi \log(1+x) + (n-\xi) \log(1-x))$, each logarithm vanishing at $x=0$. As a function of x this function is single valued in, say, the x -plane cut on $(-\infty, -1]$ and $[1, \infty)$, and the series (13) converges to it in: $|x| < 1$.

Binomial coefficients are defined by

$$\begin{aligned} \binom{\zeta}{s} &\equiv \frac{\Gamma(\zeta+1)}{s! \Gamma(\zeta+1-s)} \\ &= \frac{\zeta(\zeta-1) \cdots (\zeta-s+1)}{s!} \quad s > 0 \end{aligned}$$

when ζ is not an integer, and $\varphi_s(n, \xi)$ is seen to be a polynomial in ξ of degree s :

$$\varphi_s(n, \xi) = \frac{2^s}{s!} \xi^s + \cdots + (-1)^s \binom{n}{s} \quad (15)$$

The recurrence relation

$$\varphi_0(n, \xi) + \varphi_1(n, \xi) + \cdots + \varphi_s(n, \xi) = \varphi_s(n-1, \xi) \quad (16)$$

obtained by expanding the various factors [] in the identity

$$[(1+x)^\xi (1-x)^{n-\xi}] [(1-x)^{-1}] = [(1+x)^\xi (1-x)^{n-1-\xi}]$$

is an important one. We note also for reference that

$$\begin{aligned} \varphi_0(n, \xi) &= 1 \\ \varphi_s(n, n-\xi) &= (-1)^s \varphi_s(n, \xi) \\ \varphi_s(n, n) &= \binom{n}{s} \\ \varphi_s(n-1, n) &= 1 + \binom{n}{1} + \cdots + \binom{n}{s} \end{aligned} \quad (17)$$

valid for all n, ξ and non-negative integers s . We see, by the way, that $\varphi_e(n-1, n)$ is simply the Hamming expression (1).

The function $f_{n,\xi}(x)$ has the property that

$$f_{n,\xi}\left(\frac{x+y}{1+xy}\right) = \frac{f_{n,\xi}(x)f_{n,\xi}(y)}{(1+xy)^n} \quad (18)$$

at least if, say, given x , $|y|$ is small enough. From this and (8) for the operator L_j it is apparent that

$$L_j f_{n,\xi}(x) = \varphi_j(n, \xi) f_{n,\xi}(x) \quad (19)$$

Similarly, using (19) and (16), or directly from (10) for the operator M ,

$$\begin{aligned} M f_{n,\xi}(x) &= [L_0 + L_1 + \cdots + L_e] f_{n,\xi}(x) \\ &= \varphi_e(n-1, \xi) f_{n,\xi}(x) \end{aligned} \quad (20)$$

If ξ_β is one of the roots of the polynomial $\varphi_e(n-1, \xi)$ then (20) becomes

$$M(1+x)^{\xi_\beta}(1-x)^{n-\xi_\beta} = 0$$

If we assume for the moment for simplicity that $\varphi_e(n-1, \xi)$ has e distinct roots ξ_β , $1 \leq \beta \leq e$, then (11) has as complementary function

$$\sum_{\beta=1}^e A_\beta (1+x)^{\xi_\beta} (1-x)^{n-\xi_\beta}$$

where the A_β are e arbitrary constants.

Fortunately, the function $(1+x)^n = f_{n,n}(x)$ is also a member of the family (13); hence

$$M(1+x)^n = \varphi_e(n-1, n)(1+x)^n$$

and the function

$$\frac{(1+x)^n}{\varphi_e(n-1, n)}$$

is a particular integral of (11). [We see from (17) that $\varphi_e(n-1, n)$ does not vanish in cases of interest.] Finally, when the roots of $\varphi_e(n-1, \xi)$ are distinct, the general solution of (11) must be of the form

$$G(x) = \frac{(1+x)^n}{\varphi_e(n-1, n)} + \sum_{\beta=1}^e A_\beta (1+x)^{\xi_\beta} (1-x)^{n-\xi_\beta} \quad (21)$$

If $\varphi_e(n-1, \xi)$ has multiple roots then the general solution will contain additional terms

$$(\text{const.}) (1+x)^{\xi_\beta} (1-x)^{n-\xi_\beta} \left[\log \frac{1+x}{1-x} \right]^\mu \quad (22)$$

i.e., the μ^{th} derivative of $f_{n,\xi}(x)$ with respect to ξ , with μ any positive integer less than the multiplicity of root ξ_β .

Before applying these results to e -codes in detail, let us derive a certain modification of (21). First, we see from (17) that if n is a positive integer, then n is not one of the roots of $\varphi_e(n-1, \xi)$. If the roots ξ_β of $\varphi_e(n-1, \xi)$ are distinct and if A_β , $1 \leq \beta \leq e$, are any e numbers then a polynomial $\theta(\xi)$ of formal degree e is uniquely determined by the $e+1$ conditions:

$$\begin{aligned}\theta(\xi_\beta) &= (\xi_\beta - n)\varphi_e'(n-1, \xi_\beta)A_\beta, & 1 \leq \beta \leq e,^* \\ \theta(n) &= 1\end{aligned}\quad (23)$$

using, e.g., the Lagrange interpolation formula. It is obvious that $G(x)$, (21), may be expressed in terms of this polynomial as

$$G(x) = \frac{1}{2\pi i} \int_{\Gamma} \frac{(1+x)^\xi (1-x)^{n-\xi} \theta(\xi)}{(\xi - n)\varphi_e(n-1, \xi)} d\xi \quad (24)$$

where Γ is any simple closed contour surrounding the roots: $n, \xi_1, \xi_2, \dots, \xi_e$ of the denominator of the integrand; (the numerator is an entire function of ξ provided $x^2 \neq 1$).

Analysis a little more detailed shows that even if $\varphi_e(n-1, \xi)$ has multiple roots the general solution of (11) can be represented in the form (24), again with $\theta(\xi)$ any polynomial of formal degree e such that $\theta(n) = 1$. The e constants of integration appear as the $e+1$ parameters of $\theta(\xi)$ restricted by $\theta(n) = 1$.†

Expansion of the integrand in (24) according to (13) yields the form

$$\nu_s = \frac{1}{2\pi i} \int_{\Gamma} \frac{\varphi_s(n, \xi)\theta(\xi)}{(\xi - n)\varphi_e(n-1, \xi)} d\xi \quad s = 0, 1, 2, \dots \quad (25)$$

for the coefficients of $G(x)$, (12).

If we denote by

$$L_j G(x) \equiv G_j(x) = \sum_{s=0}^{\infty} \nu_{j,s} x^s \quad (26)$$

the result of applying the operator L_j to any solution (24) of (11), then it is straightforward that

$$G_j(x) = \frac{1}{2\pi i} \int_{\Gamma} \frac{(1+x)^\xi (1-x)^{n-\xi} \varphi_j(n, \xi)\theta(\xi)}{(\xi - n)\varphi_e(n-1, \xi)} d\xi, \quad (27)$$

* The prime denotes differentiation with respect to ξ .

† If $G(x)$ of (24) is to satisfy (11) it is sufficient that $\theta(\xi)$ be any function regular within (and on) Γ and that $\theta(n) = 1$, as may be easily verified. Since $G(x)$ depends on $\theta(\xi)$ only by way of the values of $\theta(\xi)$ at the zeros of the denominator in (24), the condition that $\theta(\xi)$ be a polynomial of formal degree e serves merely to determine $\theta(\xi)$ uniquely for a given solution $G(x)$.

and that

$$\nu_{j,s} = \frac{1}{2\pi i} \int_{\Gamma} \frac{\varphi_s(n, \xi) \varphi_j(n, \xi) \theta(\xi)}{(\xi - n) \varphi_e(n - 1, \xi)} d\xi, \quad s = 0, 1, \dots \quad (28)$$

(An interesting reciprocity $\nu_{j,s} = \nu_{s,j}$ is apparent from (28). In an e -code one has (number of j -neighbors of weight s) = (number of s -neighbors of weight j) only for $0 \leq s, j \leq e$, since $L_j G(x)$ is the generating function for j -neighbors only if $0 \leq j \leq e$.)

IV. BOUNDARY CONDITIONS

The coefficients ν_s , (25), of any solution of (11) satisfy the relation:

$$\nu_0 + \nu_1 + \dots + \nu_e = \frac{1}{2\pi i} \int_{\Gamma} \frac{\theta(\xi)}{\xi - n} d\xi = 1 \quad (29)$$

by virtue of (16) and the normalizing condition $\theta(n) = 1$.

With γ an integer such that $0 \leq \gamma \leq e$, denote by

$$G^{(\gamma)}(x) = \sum_{s=0}^{\infty} \nu_s^{(\gamma)} x^s \quad (30)$$

a solution of (11) which satisfies the e boundary conditions

$$\begin{aligned} \nu_0^{(\gamma)} &= \nu_1^{(\gamma)} = \dots = \nu_{\gamma-1}^{(\gamma)} = 0 \\ \nu_{\gamma+1}^{(\gamma)} &= \nu_{\gamma+2}^{(\gamma)} = \dots = \nu_e^{(\gamma)} = 0 \end{aligned} \quad (31)$$

We must have $\nu_{\gamma}^{(\gamma)} = 1$ in such a solution, from (29). Thus the conditions (31) are equivalent to specifying the values of $G^{(\gamma)}(x)$ and its first $e - 1$ derivatives at the ordinary point $x = 0$ of (11), so that such a solution $G^{(\gamma)}(x)$ exists and is uniquely determined.⁶

Given an e -code on n places, each n -word of B_n lies at distance e or less from exactly one code word; namely, the code word to which it belongs. In particular, the n -word $00 \cdots 0$ must lie at distance e or less from a single code word. That is to say, there is exactly one code word in the sphere of radius e centered at $00 \cdots 0$. If this code word is of weight γ , then the generating function for the given e -code can be none other than the solution $G^{(\gamma)}(x)$ of (11) defined in the preceding paragraph.

If there exists an e -code on n places in which the code word of least weight is of weight γ , then there can be derived from it an e -code on n places in which the code word of least weight is of weight γ' , where γ' is any integer satisfying $0 \leq \gamma' \leq e$. The transformation is that of choosing certain places among n and then changing the letters of each n -word of B_n in these places, 0's to 1's and 1's to 0's. (Such a transformation

corresponds to one of the operations of the orthogonal group which leaves invariant the n -cube representing B_n .) Metric properties in B_n are invariant under such a transformation, clearly, and an e -code is transformed into an e -code. Thus if there exists any e -code on n places then (11) must have $e + 1$ distinct polynomial solutions $G^{(\gamma)}(x)$, satisfying boundary conditions (31) for each case $\gamma = 0, 1, \dots, e$.

In (25) for the coefficients ν_s , move contour Γ out to a circle sufficiently large that the expansion

$$\frac{1}{(\xi - n)\varphi_e(n - 1, \xi)} = \frac{e!}{2^e \xi^{e+1}} + \frac{(\text{const.})}{\xi^{e+2}} + \dots$$

converges on Γ . Suppose that the polynomial $\theta(\xi)$, of formal degree e , is of actual degree f : $\theta(\xi) = c\xi^f + 0(\xi^{f-1})$, $c \neq 0$, where $0 \leq f \leq e$. Then the numerator of the integrand in (25) is of the form: $(2^s c \xi^{s+f}/s!) + 0(\xi^{s+f-1})$, and it is clear that

$$\begin{aligned} \nu_s &= 0 & 0 \leq s \leq e - f - 1 \\ \nu_{e-f} &= \frac{e!c}{2^f(e-f)!} \neq 0 \end{aligned}$$

Hence, if $\theta^{(\gamma)}(\xi)$ denotes the polynomial which gives $G^{(\gamma)}(x)$ in the representation (24), then $\theta^{(\gamma)}(\xi)$ must be of actual degree $e - \gamma$.

A particularly simple case is the one $\gamma = e$; the polynomial $\theta^{(e)}(\xi)$ must be of degree zero, and is determined by the normalization as $\theta^{(e)}(\xi) \equiv 1$. Thus

$$G^{(e)}(x) = \frac{1}{2\pi i} \int_{\Gamma} \frac{(1+x)^{\xi}(1-x)^{n-\xi}}{(\xi - n)\varphi_e(n - 1, \xi)} d\xi \quad (32)$$

From this we have immediately the following

Theorem: If there exists an e -code on n places then the equation $\varphi_e(n - 1, \xi) = 0$ in ξ has e distinct integer roots.

Proof: If there exists an e -code on n places, then there exists an e -code on n places in which the code word of least weight is of weight e . The solution (32) of (11) must be the generating function for this e -code; hence (32) must reduce to a polynomial of formal degree n . If $\varphi_e(n - 1, \xi)$ had multiple roots then noncancelling logarithmic terms (22) would appear in the $G^{(e)}(x)$ of (32). Thus $\varphi_e(n - 1, \xi)$ must have e distinct roots ξ_{β} , $1 \leq \beta \leq e$. Each solution $(1+x)^{\xi_{\beta}}(1-x)^{n-\xi_{\beta}}$ of the homogeneous equation appears in $G^{(e)}(x)$ with nonvanishing coefficient:

$$A_{\beta} = \frac{1}{(\xi_{\beta} - n)\varphi_e'(n - 1, \xi_{\beta})}$$

Since $G^{(e)}(x)$ must be a polynomial in x , it must be expressible as a polynomial in $1 + x$; hence each root ξ_β must be an integer.* (It is not necessary to require further that $0 \leq \xi_\beta \leq n$, since it follows easily from (14) that any real root of $\varphi_e(n - 1, \xi)$ satisfies $0 \leq \xi_\beta \leq n - 1$ provided n and e are integers such that $0 \leq e \leq n$.)

As a corollary we have that if e is odd then n must be odd. This follows from the theorem and the fact that $\frac{1}{2}(n - 1)$ is a root of $\varphi_e(n - 1, \xi)$ when e is odd, from (17).

We consider next the case $\gamma = 0$. If $00 \cdots 0$ is a code word, then its e -neighbors are the n -words of weight e . Furthermore, the n -words of weight less than e belong to the code word $00 \cdots 0$, and can be e -neighbors neither of $00 \cdots 0$ nor of any other code word. Hence it must be true that

$$G_e^{(0)}(x) = \binom{n}{e} x^e + 0(x^{e+1}) \quad (33)$$

With $G_e^{(0)}(x)$ represented in the form (27), divide the factor $\varphi_e(n, \xi)\theta^{(0)}(\xi)$ in the numerator by the denominator; the result will be

$$\varphi_e(n, \xi)\theta^{(0)}(\xi) = [(\xi - n)\varphi_e(n - 1, \xi)]q(\xi) + r(\xi) \quad (34)$$

with quotient $q(\xi)$ a polynomial of degree $e - 1$ and remainder $r(\xi)$ a polynomial of formal degree e . The term involving $q(\xi)$ obviously contributes nothing to $G_e^{(0)}(x)$ in (27), so that from (33) and arguments similar to those giving $G^{(e)}(x)$, above, $r(\xi)$ must be the constant

$$r(\xi) = \binom{n}{e} = \varphi_e(n, n)$$

From (34) we then obtain the values of $\theta^{(0)}(\xi)$ at the poles of the integrand in (24), and thus

$$G^{(0)}(x) = \frac{(1 + x)^n}{\varphi_e(n - 1, n)} + \sum_{\beta=1}^e \frac{\varphi_e(n, n)(1 + x)^{\xi_\beta}(1 - x)^{n-\xi_\beta}}{\varphi_e(n, \xi_\beta)(\xi_\beta - n)\varphi_e'(n - 1, \xi_\beta)} \quad (35)$$

Before obtaining $G^{(\gamma)}(x)$ explicitly for intermediate values of γ , we must first discuss a certain set of recursion relations holding between the coefficients ν_s of any solution of (11). These relations are

$$\sum_{s=e-\rho+1}^{e+\rho} (-1)^{e+s} k_{\rho,s} \nu_s = 0, \quad \rho = 1, 2, \dots, \quad (36)$$

* The condition of Golay for the existence of group codes, obtained by different means, is essentially that $\varphi_e(n - 1, \xi)$ have at least one root an integer. Cf.: (4) of Reference 4, in view of (16), above.

where we define $\nu_s = 0$ for $s < 0$ and where the coefficients $k_{\rho,s}$ are

$$k_{\rho,s} = \sum_{\sigma=0}^{\rho} \binom{s}{\sigma} \binom{n-s}{\rho-\sigma} \binom{\rho-1}{e-s+\sigma} \quad (37)$$

(The derivation of (36) is given in Appendix A.) Equations (36), written out, are of the form

$$\begin{aligned} k_{1,e}\nu_e - k_{1,e+1}\nu_{e+1} &= 0 \\ k_{2,e-1}\nu_{e-1} - k_{2,e}\nu_e + k_{2,e+1}\nu_{e+1} - k_{2,e+2}\nu_{e+2} &= 0 \\ &\vdots \end{aligned}$$

from which we see that (36) may be used to determine

$$\nu_{e+1}, \nu_{e+2}, \dots$$

recursively in terms of

$$\nu_e, \nu_{e-1}, \dots, \nu_0$$

We see also that if

$$\nu_e = \nu_{e-1} = \dots = \nu_{\gamma+1} = 0$$

(with γ such that $0 \leq \gamma \leq e-1$) then

$$\nu_{e+1} = \nu_{e+2} = \dots = \nu_{2e-\gamma} = 0.$$

This has the following interpretation in terms of e -codes. It is well known (and obvious) that two different code words in an e -code must be separated by distance at least $2e+1$. Hence if the code word of least weight in an e -code is of weight γ then all other code words are of weight not less than $2e+1-\gamma$. In the generating function for such a code it must be the case that not only

$$G^{(\gamma)}(x) = x^\gamma + 0(x^{e+1})$$

but in fact

$$G^{(\gamma)}(x) = x^\gamma + 0(x^{2e+1-\gamma}) \quad (38)$$

Equations (36) insure that this condition is satisfied automatically.*

As a particular case of (38), we have

$$G^{(0)}(x) = 1 + 0(x^{2e+1}).$$

We see that if we apply the operator L_γ to $G^{(0)}(x)$ there will result

$$L_\gamma G^{(0)}(x) = \varphi_\gamma(n, n)x^\gamma + 0(x^{2e+1-\gamma}) \quad (39)$$

* It is also necessary for the existence of an e -code that (36) determine $\nu_{e+1}, \nu_{e+2}, \dots$ as non-negative integers when $\nu_e, \nu_{e-1}, \dots, \nu_0$ are those of (31). This condition is discussed a little further in Appendix A.

using the differential operator form for L_γ , (9). On the other hand, the function

$$\frac{L_\gamma G^{(0)}(x)}{\varphi_\gamma(n, n)} = \frac{1}{2\pi i} \int_\Gamma \frac{(1+x)^\xi (1-x)^{n-\xi}}{(\xi-n)\varphi_e(n-1, \xi)} \left[\frac{\theta^{(0)}(\xi)\varphi_\gamma(n, \xi)}{\varphi_\gamma(n, n)} \right] d\xi \quad (40)$$

is a solution of differential equation (11), in view of the discussion following (24). From (39) we see that this function can be none other than $G^{(\gamma)}(x)$. Finally, applying L_γ to $G^{(0)}(x)$ in the form (35), we have explicitly

$$G^{(\gamma)}(x) = \frac{(1+x)^n}{\varphi_e(n-1, n)} + \frac{\varphi_e(n, n)}{\varphi_\gamma(n, n)} \sum_{\beta=1}^e \frac{\varphi_\gamma(n, \xi_\beta)(1+x)^{\xi_\beta}(1-x)^{n-\xi_\beta}}{\varphi_e(n, \xi_\beta)(\xi_\beta-n)\varphi_e'(n-1, \xi_\beta)} \quad 0 \leq \gamma \leq e. \quad (41)$$

V. EXAMPLES

The known cases where the condition of Hamming is satisfied are the following:

Case I: $e = 0, n \geq 1$

The Hamming expression (1) reduces to unity. In fact,

$$\varphi_0(n-1, \xi) \equiv 1,$$

and the condition that all roots be integers is vacuous. The generating function for code words is (uniquely):

$$G^{(0)}(x) = \frac{(1+x)^n}{\varphi_0(n-1, n)} = (1+x)^n$$

Each n -word of B_n is a detection region and thus a code word. There is no error correction.

Case II: $e \geq 1, n = e$

The Hamming expression becomes the sum of all the terms in the binomial expansion of $(1+1)^n$. The "codes" in this class consist of a single code word surrounded by its detection region consisting of the sphere B_n of radius n . No signalling is possible, of course, but our methods still apply.

From the representation

$$\begin{aligned} \varphi_s(n, \xi) &= \frac{1}{2\pi i} \int_C \frac{(1+x)^\xi (1-x)^{n-\xi}}{x^{s+1}} dx \\ &= \frac{1}{2\pi i} \int_C \frac{(1+2v)^\xi}{v^{s+1}(1+v)^{n-s+1}} dv \end{aligned} \quad (42)$$

(valid for all n, ξ) we have immediately

$$\varphi_n(n-1, \xi) = 2^n \binom{\xi}{n} = \frac{2^n}{n!} \xi(\xi-1) \cdots (\xi-n+1)$$

and the roots are $0, 1, \cdots, n-1$. The generating function $G^{(e)}(x)$ of (32) becomes*

$$\begin{aligned} G^{(n)}(x) &= \frac{n!}{2^n} \cdot \frac{1}{2\pi i} \int_{\Gamma} \frac{(1+x)^\xi (1-x)^{n-\xi}}{(\xi)_{n+1}} d\xi \\ &= \frac{n!}{2^n} \sum_{\xi=0}^n \frac{(-1)^{n-\xi}}{\xi!(n-\xi)!} (1+x)^\xi (1-x)^{n-\xi} \\ &= \frac{1}{2^n} [(1+x) - (1-x)]^n = x^n \end{aligned}$$

as one might expect. The explicit form for $\varphi_n(n, \xi)$ is somewhat complicated, but for ξ an integer it follows immediately from definition (13) that

$$\varphi_n(n, \xi) = (-1)^{n-\xi} \quad \xi = 0, 1, \cdots, n$$

From (35), then,

$$\begin{aligned} G^{(0)}(x) &= \frac{1}{2^n} \sum_{\xi=0}^n \binom{n}{\xi} (1+x)^\xi (1-x)^{n-\xi} \\ &= \frac{1}{2^n} [(1+x) + (1-x)]^n = 1 \end{aligned}$$

which, again, is not surprising. The details for other values of γ seem to be more tedious, although one expects (41) to yield $G^{(\gamma)}(x) = x^\gamma$.

Case III: $e \geq 1, n = 2e + 1$

The Hamming expression in this case:

$$1 + \binom{2e+1}{1} + \cdots + \binom{2e+1}{e} = 2^{2e}$$

consists of the first half of the terms in the binomial expansion of $(1+1)^{2e+1}$. The code words in a code of this class are any two n -words separated by distance n (i.e., two vertices at opposite corners of the n -cube). The group codes in this class are the "majority rule" codes.† From (42) we have (using the substitution $y = 4v + 4v^2$)

* $(\zeta)_s \equiv s! \binom{\zeta}{s}$ denotes the descending factorial.

† The two code words in such a code are $00 \cdots 0$ and $11 \cdots 1$. An n -word belongs to $00 \cdots 0$ if it contains more 0's than 1's, and to $11 \cdots 1$ if it contains more 1's than 0's.

$$\varphi_s(n, \xi) = \frac{2^n}{2\pi i} \int_c \frac{(1+y)^{\frac{1}{2}(\xi-1)} dy}{[(1+y)^{\frac{1}{2}} - 1]^{s+1} [(1+y)^{\frac{1}{2}} + 1]^{n-s+1}}, \quad (43)$$

and, without difficulty,

$$\varphi_e(2e, \xi) = 2^{2e} \binom{\frac{1}{2}(\xi-1)}{e} = \frac{2^e}{e!} (\xi-1)(\xi-3) \cdots (\xi-2e+1)$$

The roots are 1, 3, $\dots$, $2e-1$, and from (32):

$$\begin{aligned} G^{(e)}(x) &= \frac{e!}{2^e} \cdot \frac{1}{2\pi i} \int_{\Gamma} \frac{(1+x)^{\xi}(1-x)^{2e+1-\xi} d\xi}{(\xi-2e-1)(\xi-1)(\xi-3) \cdots (\xi-2e+1)} \\ &= 2^{-2e} (1+x)(1+x)^2 - (1-x)^2]^e = x^e + x^{e+1} \end{aligned}$$

In the case $\gamma = 0$ we need the result

$$\varphi_e(2e+1, \xi) = 2^{2e+1} \left[\binom{\frac{1}{2}\xi}{e+1} - \binom{\frac{1}{2}(\xi-1)}{e+1} \right]$$

from (43). It is then tedious but straightforward to obtain from (35)

$$\begin{aligned} G^{(0)}(x) &= \frac{1}{2^{2e}} \sum_{r=0}^e \binom{2e+1}{2r+1} (1+x)^{2r+1} (1-x)^{2e-2r} \\ &= 2^{-2e-1} \{ [(1-x) + (1+x)]^{2e+1} - [(1-x) - (1+x)]^{2e+1} \} \\ &= 1 + x^{2e+1} \end{aligned}$$

One expects to get

$$G^{(\gamma)}(x) = x^{\gamma} + x^{2e+1-\gamma}$$

from (41), but verification appears to be complicated.

Case IV: $e = 1, n = 2^t - 1$ ($t = 3, 4, \dots$)

The single error correcting codes of Hamming¹ are included here. (The examples for $t = 1$, resp. $t = 2$, appear under Case II, resp. Case III, above.) Since n is always odd the condition that $\varphi_1(n-1, \xi) = 2\xi - n + 1$ have an integer root is automatically satisfied. For $\gamma = 1$ the generating function is

$$\begin{aligned} G^{(1)}(x) &= \frac{1}{2\pi i} \int_{\Gamma} \frac{(1-x)^{\xi}(1-x)^{n-\xi}}{(\xi-n)(2\xi-n+1)} d\xi \\ &= \frac{(1+x)^n - (1+x)^{\frac{1}{2}(n-1)}(1-x)^{\frac{1}{2}(n+1)}}{1+n} \end{aligned}$$

from which we have

$$\nu_s^{(1)} = \frac{1}{1+n} \left\{ \binom{n}{s} - (-1)^{\frac{1}{2}s} \binom{\frac{1}{2}(n-1)}{\frac{1}{2}s} \right\} \quad s \text{ even}$$

$$\nu_s^{(1)} = \frac{1}{1+n} \left\{ \binom{n}{s} - (-1)^{\frac{1}{2}(s+1)} \binom{\frac{1}{2}(n-1)}{\frac{1}{2}(s-1)} \right\} \quad s \text{ odd}$$

For $\gamma = 0$, Eq. (35) works out as

$$G^{(0)}(x) = \frac{(1+x)^n + n(1+x)^{\frac{1}{2}(n-1)}(1-x)^{\frac{1}{2}(n+1)}}{1+n}$$

so that

$$\begin{aligned} \nu_s^{(0)} &= \frac{1}{1+n} \left\{ \binom{n}{s} + n(-1)^{\frac{1}{2}s} \binom{\frac{1}{2}(n-1)}{\frac{1}{2}s} \right\} \quad s \text{ even} \\ \nu_s^{(0)} &= \frac{1}{1+n} \left\{ \binom{n}{s} + n(-1)^{\frac{1}{2}(s+1)} \binom{\frac{1}{2}(n-1)}{\frac{1}{2}(s-1)} \right\} \quad s \text{ odd} \end{aligned}$$

Case V: $e = 2, n = 90$

The double error correcting codes for $n = 2, 5$ are covered by Cases II, III, respectively. The discovery that

$$1 + 90 + \frac{1}{2}(90)(89) = 2^{12}$$

is due to Golay.⁷ We have

$$2\varphi_2(n-1, \xi) = (2\xi - n + 1)^2 - (n-1)$$

with roots

$$\frac{1}{2}[n-1 \pm (n-1)^{\frac{1}{2}}]$$

Since these roots are not integers when $n = 90$, there can be no 2-code for $n = 90$.* H. S. Shapiro has shown (in unpublished work) that the Hamming condition for $e = 2$ is satisfied only in the cases $n = 2, 5, 90$, so that the only nontrivial 2-codes are those equivalent to the majority rule code on 5 places.

Case VI: $e = 3, n = 23$

Golay⁷ finds:

$$1 + 23 + \frac{1}{2}(23)(22) + (23)(22)(21)/6 = 2^{11}$$

and gives explicitly a 3-code on 23 places of group type. We have

$$6\varphi_3(n-1, \xi) = (2\xi - n + 1)[(2\xi - n + 1)^2 - (3n - 5)]$$

and when $n = 23$ we verify that the roots are the integers 7, 11, 15. Computations by the author show that for $n < 10^{10}$ the Hamming condition for $e = 3$ holds only when $n = 3, 7, 23$.

* This settles a question raised by Golay, who shows that there is no code of group type in this case, but not that there is no code at all.

For $e = 4$ we have

$$24\varphi_4(n-1, \xi) = [(2\xi - n + 1)^2 - (3n - 7)]^2 - (6n^2 - 30n + 40)$$

For $n = 4, 9$ this reduces to the forms given under Cases II, III. Preliminary calculations by the author shows that any other solutions of the Hamming condition for $e = 4$ must be such that $n > 10^{10}$, so that the question of the existence of 4-codes (other than the majority rule code) is somewhat academic.

Computations of Mrs. G. Rowe of the Mathematical Research Department show that Cases I-VI cover all cases of the Hamming condition being satisfied in the range

$$0 \leq e \leq n, \quad 1 \leq n \leq 150$$

APPENDIX A

From (13) we have

$$\left(\frac{1-x}{1+x}\right)^{n-\xi} = \frac{1}{(1+x)^n} \sum_{s=0}^{\infty} \varphi_s(n, \xi) x^s \quad (\text{A1})$$

Applying the operator $D = -(1+x)^2 d/dx$ to both sides of (A1) ρ times, there results

$$2^\rho (n-\xi)_\rho \left(\frac{1-x}{1+x}\right)^{n-\xi-\rho} = \sum_{s=0}^{\infty} \varphi_s(n, \xi) D^\rho [x^s (1+x)^{-n}] \quad (\text{A2})$$

The substitution $v = (1+x)^{-1}$ reduces D to d/dv , so that

$$\begin{aligned} D^\rho x^s (1+x)^{-n} &= \frac{d^\rho}{dv^\rho} (1-v)^s v^{n-s} \\ &= \sum_{\sigma=0}^{\rho} \binom{\rho}{\sigma} (-1)^\sigma (s)_\sigma (1-v)^{s-\sigma} (n-s)_{\rho-\sigma} v^{n-s-\rho+\sigma} \\ &= \rho! \sum_{\sigma=0}^{\rho} (-1)^\sigma \binom{s}{\sigma} \binom{n-s}{\rho-\sigma} x^{s-\sigma} (1+x)^{\rho-n} \end{aligned}$$

using Leibnitz's rule. We substitute this into Eq. (A2), multiply both sides of the result by

$$(1+x)^{n-\rho} (1-x)^{\rho-\tau} / \rho!$$

(with τ arbitrary), and then equate coefficients of x^t on both sides; there obtains

$$2^\rho \binom{n-\xi}{\rho} \varphi_t(n-\tau, \xi) = \sum_{s=0}^{t+\rho} (-1)^{t+s} \kappa_{\rho,s}(n, \tau; t) \varphi_s(n, \xi) \quad (\text{A3})$$

valid for all n, ξ, τ and all non-negative integers ρ, t , where

$$\kappa_{\rho,s}(n, \tau; t) = \sum_{\sigma=0}^{\rho} \binom{s}{\sigma} \binom{n-s}{\rho-\sigma} \binom{\rho-\tau}{t-s+\sigma} \quad (\text{A4})$$

The coefficients $\kappa_{\rho,s}(n, \tau; t)$ vanish unless $s \leq t + \rho$; if n and $\rho - \tau$ are non-negative integers then the coefficients $\kappa_{\rho,s}(n, \tau; t)$ are positive integers provided $t - \rho + \tau \leq s$ and vanish otherwise. In particular, (setting $\tau = 1, t = e$),

$$2^{\rho} \binom{n-\xi}{\rho} \varphi_e(n-1, \xi) = \sum_{s=e-\rho+1}^{e+\rho} (-1)^{e+s} k_{\rho,s} \varphi_s(n, \xi), \rho = 1, 2, \dots, \quad (\text{A5})$$

where we define $\varphi_s(n, \xi) \equiv 0$ for $s < 0$; the $k_{\rho,s} = \kappa_{\rho,s}(n, 1; e)$ are those of (37) of the text. If we multiply ν_s of (25) by $(-1)^{e+s} k_{\rho,s}$ and sum on s there results (36), since the left hand member of (A5) is a polynomial multiple of the denominator of the integrand in (25).

If the code word of least weight in an e -code is of weight γ , then the first nontrivial one of the (37) is the one for $\rho = e + 1 - \gamma$, and it gives (since $\nu_{\gamma}^{(\gamma)} = 1$)

$$\begin{aligned} \nu_{2e+1-\gamma}^{(\gamma)} &= \frac{k_{e+1-\gamma,\gamma}}{k_{e+1-\gamma,2e+1-\gamma}} \\ &= \frac{(n-\gamma)(n-\gamma-1) \cdots (n-e)}{(2e+1-\gamma)(2e-\gamma) \cdots (e+1)} \end{aligned}$$

A necessary condition for the existence of an e -code on n places is that this expression be a non-negative integer in each case $\gamma = 0, 1, \dots, e$. It is not clear, however, that this condition is independent of the one set forth in the theorem of Section IV.

APPENDIX B

We give here a relation due to K. M. Case⁸ which shows that the statement of our main result as it appears in the Abstract heading this article agrees with the theorem proved in Section IV.

In the defining relation

$$(1+x)^r(1-x)^{n-r} = \sum_{s=0}^{\infty} x^s \varphi_s(n, r) \quad (\text{B1})$$

for the coefficients $\varphi_s(n, r)$ assume that n and r are integers, multiply both sides by $(-1)^r \binom{n}{r} y^r$, and sum on $r, 0 \leq r \leq n$. The result is

$$[(1-x) - y(1+x)]^n = \sum_{r=0}^n \sum_{s=0}^n (-1)^r \binom{n}{r} y^r x^s \varphi_s(n, r) \quad (\text{B2})$$

Rearrange the left hand member and re-expand it, to get

$$\begin{aligned} [(1-x) - y(1+x)]^n &= [(1-y) - x(1+y)]^n \\ &= \sum_{s=0}^n (-1)^s \binom{n}{s} x^s (1+y)^s (1-y)^{n-s} \quad (\text{B3}) \\ &= \sum_{s=0}^n \sum_{r=0}^n (-1)^s \binom{n}{s} x^s y^r \varphi_r(n, s) \end{aligned}$$

Comparing coefficients of $x^s y^r$ in (B2) and (B3), we have, finally,

$$(-1)^r \binom{n}{r} \varphi_s(n, r) = (-1)^s \binom{n}{s} \varphi_r(n, s) \quad (n, r, s \text{ integers}), \quad (\text{B4})$$

or, changing notation slightly,

$$\varphi_\xi(n-1, e) = (-1)^{\xi-e} \frac{\binom{n-1}{\xi}}{\binom{n-1}{e}} \varphi_e(n-1, \xi) \quad (\text{B5})$$

(with n, e, ξ integers and $0 \leq e, \xi \leq n-1$). Thus if $\varphi_e(n-1, \xi)$ vanishes for e different integers ξ then so must $\varphi_\xi(n-1, e)$, at least when $e < n$. But $\varphi_\xi(n-1, e)$ is the coefficient of x^ξ in $(1+x)^e (1-x)^{n-1-e}$ when this is written out as a polynomial in x , by definition.

REFERENCES

1. R. W. Hamming, B.S.T.J., **29**, p. 147, 1950.
2. C. E. Shannon, B.S.T.J., **27**, p. 379, 1948.
3. P. Elias, Trans. I.R.E., **PGIT-4**, p. 29, 1954.
4. M. J. E. Golay, Trans. I.R.E., **PGIT-4**, p. 23, 1954.
5. D. Slepian, B.S.T.J., **35**, p. 203, 1956.
6. E. L. Ince, *Ordinary Differential Equations* (Dover), Ch. V, XV.
7. M. J. E. Golay, Proc. I.R.E., **37**, p. 637, 1949.
8. K. M. Case, Phys. Rev., **97**, p. 810, 1955.

Selecting the Best One of Several Binomial Populations

By MILTON SOBEL and MARILYN J. HUYETT

(Manuscript received July 17, 1956)

Tables have been prepared for use in any experiment designed to select that particular one of k binomial processes or populations with the highest (long time) yield or the highest probability of success. Before experimentation the experimenter chooses two constants d^ and P^* ($0 < d^* \leq 1$; $0 \leq P^* < 1$) and specifies that he would like to guarantee a probability of at least P^* of a correct selection whenever the true difference between the long-time yields associated with the best and the second best processes is at least d^* . The tables show the smallest number of units required per process to be put on test to satisfy this specification. Separate tables are given for $k = 2, 3, 4$ and 10 . Each table gives the result for $d^* = 0.05$ (0.05) 0.50 and for $P^* = 0.50, 0.60, 0.70, 0.75, 0.80, 0.85, 0.90, 0.95$, and 0.99 . For values of d^* and P^* not considered in the tables, graphs are given on which interpolation can be carried out. Graphs have also been constructed to make possible an interpolation or extrapolation for other values of k . An alternative specification is given for use when the experimenter has some a priori knowledge of the processes and their probabilities of success. This specification is then compared with the original specification. Applications of these tables to different types of problems are considered.*

INTRODUCTION AND SUMMARY

A frequently encountered problem is that of selecting the "best" one of k ($k \geq 2$) processes or populations on the basis of the same number n of observations from each process. We shall assume that the given processes are all binomial or "go — no go" processes and that the best process is the one with the highest probability of obtaining a "success" on a single observation. We shall consider a single sample or nonsequential procedure which means that the common number n of observations from each process is to be determined before experimentation starts. The corresponding sequential problem is being investigated.¹

Briefly, the technique employed here is to let the experimenter decide how "close" the best and second best processes can be before he is willing to relax his control on the probability of a correct selection. The selection of a best process will, of course, be made on the basis of the largest observed frequency of "success"; the only remaining problem is to determine the value of n . Tables and graphs which cover almost all practical problems in this framework are given for determining the required value of n . In particular, tables and graphs are given for $k = 2, 3, 4$ and 10 . Graphs are also given to approximate the result for any value of k up to 100 .

This problem arises in many widely different fields of endeavor; we shall briefly consider two industrial applications. One application of the binomial problem is to comparative yield studies. Here success corresponds to the making of a good unit, and the goal is to select the process with the highest (long-time) yield. Another application of the binomial problem is to comparative life testing studies. In this case the experimenter selects a fixed time T and defines the best process as the one for which the probability of any one unit surviving this time T is highest. Then, of course, a successful unit is one which survives the time T . In treating this as a binomial we are discarding the information contained in the exact times of failure. In many cases the times of failure are either unknown or very inexact; in other cases it is not known how to utilize the knowledge of the exact times of failure. Hence, it would be valuable to know the results for the more basic binomial problem. The time T is considered fixed throughout; its value is determined by non-statistical considerations. The specification and the final decision of the experimenter all refer to this predetermined time T . It should be noted that the experimenter cannot use information obtained from the continuation of the test beyond time T since the best process for T is not necessarily the best process for a longer time, say $10T$. This binomial type of analysis has the advantage that it does not assume any particular form of the life distribution. In particular, the assumption of exponential life is avoided.

The presence of *à priori* information changes the number of observations required. An alternative specification is given which is justified by certain *à priori* information based on past experimentation. The amount of saving is briefly examined. This area of utilizing *à priori* information to reduce the number of observations required should be investigated further.

The treatment of the problem in this paper is based on the assumption that, after experimentation is carried out, the experimenter must

choose one of the k processes and assert that it is best. If he allows himself the possibility of hedging and asserting that he needs further experimentation, then the problem changes and the tables of this paper are not appropriate.

- The following additional assumptions will be made:
- 1. Observations from the same or different processes are independent.
 - 2. Observations from the same process have a common fixed probability of "success".
 - 3. There is no chance of error in determining whether a success or a failure has occurred.

The assumption of a common probability fixed once and for all for each process is one that should be checked carefully in any practical application of the results in this paper. Roughly speaking, this assumption states that each of the processes is in a state of statistical control as far as the probability of success is concerned.

We shall consider only the case in which the same number n of observations are taken from each process. This is certainly reasonable for a single sample procedure if no *a priori* information is assumed.

STATISTICAL FORMULATION OF THE PROBLEM

Each of k given binomial populations Π_i is associated with a fixed probability of success p_i where $0 \leq p_i \leq 1$ ($i = 1, 2, \dots, k$). For example, in the yield problem p_i is the long-time yield for process Π_i or the probability that any one unit from Π_i is a good one. Let the ordered values of the p_i be denoted by

$$p_{[1]} \geq p_{[2]} \geq \dots \geq p_{[k]} \tag{1}$$

No *a priori* information is assumed about the values of the p_i or about the correspondence between the ordered $p_{[i]}$ and the k identifiable populations Π_i . In particular, we have no idea before experimentation starts whether $p_{[1]}$ is associated with $\Pi_1, \Pi_2, \dots$, or Π_k .

The problem is to select the population associated with $p_{[1]}$ on the basis of n observations from each population. If there are t ties for first place, say

$$p_{[1]} = p_{[2]} = \dots = p_{[t]} > p_{[t+1]} \qquad (t < k) \tag{2}$$

then we shall certainly be content with the selection of any one of the associated t populations as the best one.

As an index of the true difference (or distance) between the best and second best populations we introduce the symbol

$$d = p_{[1]} - p_{[2]} \tag{3}$$

It is assumed that if the difference d between the best and second best populations is small enough, then the error involved in wrongly selecting the second best process as the best one is an error of little or no consequence. The experimenter is therefore asked to specify two quantities which will determine the number n of observations he is required to take from each process.

Specification: He specifies the smallest value d^* ($0 < d^* \leq 1$) of d for which it would be economically desirable to make the correct selection. He also specifies (4)
 a probability P^* ($0 \leq P^* < 1$) of making a correct selection that he would like to guarantee whenever the true difference $d \geq d^*$.

Letting $P_{cs} = P_{cs}(p_{[1]}, \dots, p_{[k]})$ denote the probability of a correct selection we can now rewrite the specification that the experimenter wants to satisfy in the simple form

$$P_{cs} \geq P^* \quad \text{for } d \geq d^* \quad (5)$$

[The word "specification" will be used below to denote the specified pair of constants (d^*, P^*) as well as the condition (5); it will be clear from the text which is meant.] Since the final selection is to be made on the basis of the observed frequency of success, the essential problem is to find the number n of observations required per process to satisfy the specification (5).

The possibility that d may be less than d^* is not being overlooked. The region $d < d^*$ is being regarded as a zone of indifference in the sense that if $d < d^*$, then we do not care which process is selected as best so long as its p -value is within d^* of the highest p -value $p_{[1]}$. For values of $P^* \leq 1/k$ no tables are needed since a probability of $1/k$ can be attained by chance alone.

Some comments on the above approach and on a possible modification have been placed in Appendix I in order to preserve the continuity of the paper.

CONFIDENCE STATEMENT

After the experiment is completed and the selection of a best process is made, the experimenter can make a confidence statement with confidence level P^* . Let p_s denote the true p -value of the selected population and let p_u denote the maximum true p -value over all unselected popula-

tions. Then the confidence statement, consisting of two sets of inequalities

$$\left\{ \begin{array}{l} p_{[1]} - d^* \leq p_s \leq p_{[1]} \\ p_{[2]} \leq p_u \leq p_{[2]} + d^* \end{array} \right\}$$

or

$$\left\{ \begin{array}{l} 0 \leq p_{[1]} - p_s \leq d^* \\ 0 \leq p_u - p_{[2]} \leq d^* \end{array} \right\}$$

has confidence level P^* . It should be noted that the above confidence statement is not a statement about the value of any p but is a statement about the correctness of the selection made.

LEAST FAVORABLE CONFIGURATION

The main idea used in the construction of the tables was that of a least favorable configuration. Before defining this concept we shall define the set of configurations

$$p_{[1]} - d = p_{[2]} = p_{[3]} = \cdots = p_{[k]}$$

(6)

obtained by letting d in (6) vary over the closed interval $(d^*, 1)$ as the *Less-Favorable* set of configurations. It is intuitively clear and will be rigorously shown in Appendix II that if our procedure satisfies the specification for any true configuration (6) with $d = d^0$ and $p_{[1]} = p_{[1]}^0$, then it will also satisfy the specification when

$$p_{[1]}^0 - d^0 \geq p_{[2]} \geq p_{[3]} \geq \cdots \geq p_{[k]}$$

(7)

Of course, we shall be interested particularly in the case in which d equals the specified value d^* . If $d = d^*$ is fixed in (6), then (6) specifies the differences between the p -values, but the "location" of the set is still not specified. We shall use $p_{[1]}$ to locate the set of p -values. The probability P_{cs} of a correct selection for configurations like (6) with $d = d^*$ depends not only on d^* , n and k but also on the location $p_{[1]}$ of the largest p -value (except for the special case $k = 2$ and $n = 1$). [In the corresponding problem for selecting the largest population mean of k independent *Normal* distributions with unit variance,² this probability P_{cs} depends *only on the differences* and, hence, only on d in the configuration corresponding to (6)].

When (6) holds with any *fixed* value of d , the probability P_{cs} (for any fixed n) may be regarded as a function of $p_{[1]}$ where $d \leq p_{[1]} \leq 1$). This function is continuous and bounded over a closed interval and therefore assumes its minimum value at some point $p_{[1]}(d) = p_{[1]}(d;n)$ in the closed interval $(d,1)$. Fig. 1(b) gives the value of $p_{[1]}(d)$ as a function of d for $k = 3$ and for $n = 1, 2, 4, 10$ and ∞ . For any particular value

of n and for $d = d^*$ we shall be particularly interested in the value $p_{[1]}^L = p_{[1]}(d^*; n)$ since this (as shown in Appendix II) gives the smallest probability P_{CS}^L of a correct selection for all the configurations included in the statement of the experimenter's specification. This particular configuration (6) with $d = d^*$ and $p_{[1]} = p_{[1]}^L$ (which depends on n) is called the *Least-Favorable Configuration*.

Although the least favorable configuration depends on n , it has been empirically found that for $n \geq 10$ (and in some cases for $n \geq 4$) the least favorable configuration is approximately given by $p_{[1]}^L = \frac{1}{2}(1 + d^*)$ in which the two values, $p_{[1]}^L$ and $p_{[2]}^L = p_{[1]}^L - d^*$ are symmetric about $\frac{1}{2}$. This symmetric configuration clearly does not depend on n . Fig. 1(b) shows that as $n \rightarrow \infty$ the least favorable configuration approaches this symmetric configuration (i.e., the straight line marked $n = \infty$) quite rapidly for any value of d . In Appendix III it is proved that the symmetric configuration is least favorable as $n \rightarrow \infty$. Fig. 1(a) shows for $k = 3$, $n = 10$, and any value of d the error in P_{CS} which arises as a result of using the symmetric configuration instead of the true least favorable configuration.

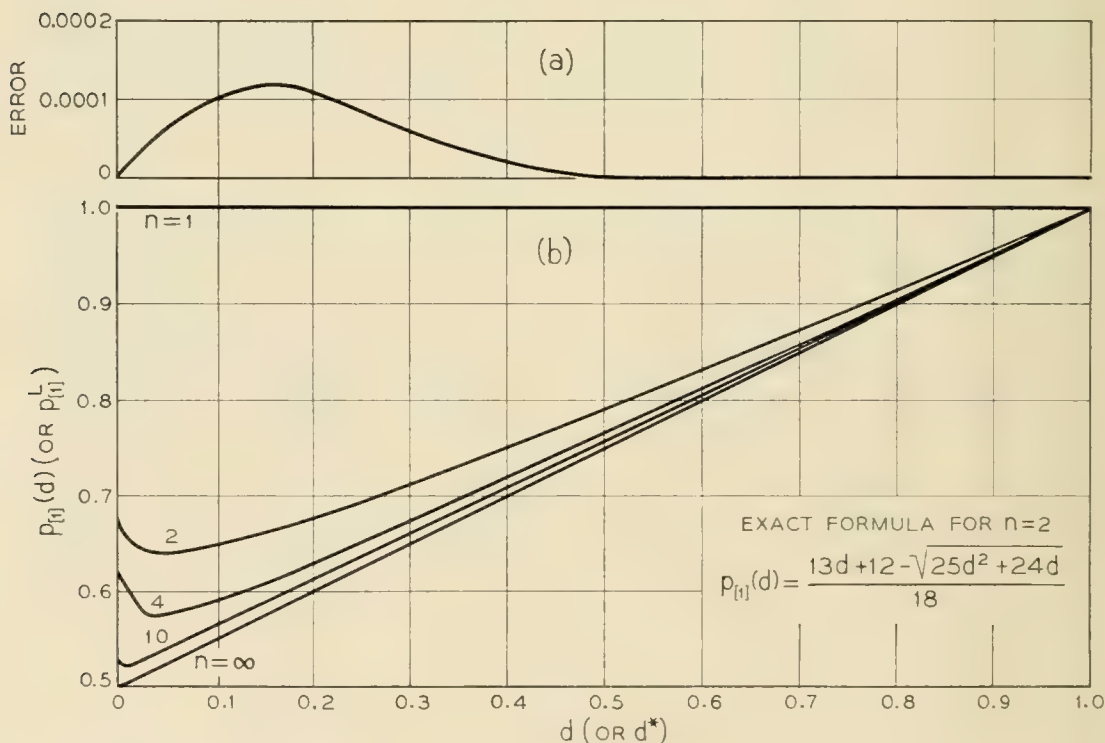


Fig. 1 — (a) Error in P_{CS} as a result of using the symmetric configuration instead of the least favorable configuration for $k = 3$, $n = 10$, and any common true difference d . (b) Least favorable value $p_{[1]}(d)$ of $p_{[1]}$ as a function of the common true difference $d = p_{[1]} - p_{[2]}$, $i \geq 2$, for $k = 3$ and selected values of n . (for $d = d^*$, $p_{[1]}(d) = p_{[1]}^L$)

CONSTRUCTION OF THE TABLES

Consider any fixed value of d^* . For each of a set of increasing values of n the minimum probability P_{cs}^L of a correct selection for $d \geq d^*$ (i.e., the probability for the least favorable configuration) was computed. These calculations were then inverted to find the smallest n for which the P_{cs}^L is greater than or equal to the specified value P^* . Tables I through IV give the smallest value of n for $k = 2, 3, 4$, and 10 , for $d^* = 0.05$ (0.05) 0.50 , and for selected values of P^* . Graphs corresponding to these tables are given in Figs. 2 through 5.

For small values of n (say, $n < 10$) it was necessary to approximate $p_{[1]}^L$ by calculating the P_{cs} *exactly* for several values of $p_{[1]}$ and proceeding in the direction of the minimum probability P_{cs}^L . For the special case $n = 2$ and $k = 3$ an explicit formula for $p_{[1]}^L$ is given on Fig. 1.

For large values of n (say, $n > 10$) the P_{cs}^L was calculated by assuming the symmetric configuration. Here it was necessary to make use of the normal approximation to the binomial. Fortunately the appropriate table needed in this normal approximation is already published.² The proof that this table is appropriate is given in Appendix III. The resulting value of n is given by

$$n = \frac{B}{d^{*2}} (1 - d^{*2}) \cong \frac{B}{d^{*2}} \tag{8}$$

where the constant B , depending on P^* and k , is equal to $\frac{1}{4}C^2$ and C is the entry in the appropriate column of Table I of R. E. Bechhofer's paper.² A short table of B values, Table V (see page 550), is included in this paper to make it self-contained.

The middle expression in (8) will be referred to as the normal approximation and the right hand expression in (8) will be referred to as the "straight line" approximation. In many cases it has been *empirically* found that these two expressions give close lower and upper bounds to the true value. Thus by noting the curves drawn in Figs. 4 and 6 for $k = 4$, $P^* \geq 0.75$ it appears that for all values of d^* the true P_{cs}^L is between the normal approximation and the straight line approximation. Assuming this to be so, it follows that for $k = 4$, $P^* \geq 0.75$ the required value of n satisfies the inequalities

$$\left[\frac{B}{d^{*2}} (1 - d^{*2}) \right] \leq n \leq \left[\frac{B}{d^{*2}} \right] \tag{9}$$

where $[x]$ denotes the smallest integer greater than or equal to the enclosed quantity x . This result (9) is empirical and not based on any mathematically proven inequalities. It is used here only to estimate the

TABLE I — NUMBER OF UNITS REQUIRED PER PROCESS TO GUARANTEE A PROBABILITY OF P^* OF SELECTING THE BEST OF k BINOMIAL PROCESSES WHEN THE TRUE DIFFERENCE $p_{[1]} - p_{[2]}$ IS AT LEAST d^* . ($k = 2$)

The three values in each group are: (1) Normal approximation, (2) Straight line approximation, and (3) Smallest integer required.

d^*	P^*							
	0.50	0.60	0.75	0.80	0.85	0.90	0.95	0.99
0.05	0	12.81	90.77	141.30	214.29	327.66	539.77	1079.70
	0	12.84	90.99	141.66	214.83	328.48	541.12	1082.41
	0	14	92	142	215	329	541	1082
0.10	0	3.18	22.52	35.06	53.17	81.30	133.93	267.90
	0	3.21	22.75—	35.41	53.71	82.12	135.28	270.60
	0	4	23	36	54	83	135	270
0.15	0	1.39	9.88	15.39	23.33	35.68	58.78	117.57
	0	1.43	10.11	15.74	23.87	36.50—	60.12	120.27
	0	2	11	16	24	37	60	120
0.20	0	0.77	5.46	8.50—	12.89	19.71	32.47	64.94
	0	0.80	5.69	8.85+	13.43	20.53	33.82	67.65+
	0	1	6	9	14	21	34	67
0.25	0	0.48	3.41	5.31	8.06	12.32	20.29	40.59
	0	0.51	3.64	5.67	8.59	13.14	21.64	43.30
	0	1	4	6	9	14	22	42
0.30	0	0.32	2.30	3.58	5.43	8.30	13.68	27.36
	0	0.36	2.53	3.93	5.97	9.12	15.03	30.07
	0	1	3	4	6	9	15	29
0.35	0	0.23	1.63	2.54	3.85—	5.88	9.69	19.38
	0	0.26	1.86	2.89	4.38	6.70	11.04	22.09
	0	1	2	3	5	7	11	21
0.40	0	0.17	1.19	1.86	2.82	4.31	7.10	14.21
	0	0.20	1.42	2.21	3.36	5.13	8.46	16.91
	0	1	2	3	4	5	9	16
0.45	0	0.13	0.90	1.39	2.11	3.23	5.33	10.65+
	0	0.16	1.12	1.75—	2.65+	4.06	6.68	13.36
	0	1	2	2	3	4	7	13
0.50	0	0.10	0.68	1.06	1.61	2.46	4.06	8.12
	0	0.13	0.91	1.42	2.15—	3.28	5.41	10.82
	0	1	1	2	3	4	5	10

TABLE II — NUMBER OF UNITS REQUIRED PER PROCESS TO GUARAN-
TEE A PROBABILITY OF P^* OF SELECTING THE BEST OF k BINOMIAL
PROCESSES WHEN THE TRUE DIFFERENCE $p_{[1]} - p_{[2]}$ IS AT LEAST
 d^* . ($k = 3$)

The three values in each group are: (1) Normal approximation, (2)
Straight line approximation, and (3) Smallest integer required.

d^*	P^*							
	0.50	0.60	0.75	0.80	0.85	0.90	0.95	0.99
0.05	30.89	78.16	205.06	272.36	363.06	496.14	732.63	1305.21
	30.97	78.36	205.58	273.04	363.97	497.38	734.46	1308.49
	31	79	206	273	364	498	735	1308
0.10	7.66	19.39	50.88	67.58	90.08	123.10	181.78	323.85+
	7.74	19.59	51.39	68.26	90.99	124.34	183.62	327.12
	8	20	52	69	91	125	184	327
0.15	3.36	8.51	22.33	29.66	39.53	54.02	79.77	142.12
	3.44	8.71	22.84	30.34	40.44	55.26	81.61	145.39
	4	9	23	31	41	55	82	145
0.20	1.86	4.70	12.33	16.38	21.84	29.84	44.07	78.51
	1.94	4.90	12.85-	17.07	22.75-	31.09	45.90	81.78
	3	5	13	17	23	31	46	81
0.25	1.16	2.94	7.71	10.24	13.65-	18.65+	27.54	49.07
	1.24	3.13	8.22	10.92	14.56	19.90	29.38	52.34
	2	4	9	11	15	20	29	52
0.30	0.78	1.98	5.20	6.90	9.20	12.57	18.57	33.08
	0.86	2.18	5.71	7.58	10.11	13.82	20.40	36.35-
	2	3	6	8	10	14	20	35
0.35	0.55+	1.40	3.68	4.89	6.52	8.91	13.15+	23.43
	0.63	1.60	4.20	5.57	7.43	10.15	14.99	26.70
	2	2	5	6	8	10	15	26
0.40	0.41	1.03	2.70	3.58	4.78	6.53	9.64	17.17
	0.48	1.22	3.21	4.27	5.69	7.77	11.48	20.45-
	1	2	4	5	6	8	11	20
0.45	0.30	0.77	2.02	2.69	3.58	4.90	7.23	12.88
	0.38	0.97	2.54	3.37	4.49	6.14	9.07	16.15+
	1	2	3	4	5	6	9	15
0.50	0.23	0.59	1.54	2.05-	2.73	3.73	5.51	9.81
	0.31	0.78	2.06	2.73	3.64	4.97	7.34	13.08
	1	2	3	3	4	5	7	12

TABLE III — NUMBER OF UNITS REQUIRED PER PROCESS TO GUARANTEE A PROBABILITY OF P^* OF SELECTING THE BEST OF k BINOMIAL PROCESSES WHEN THE TRUE DIFFERENCE $p_{[1]} - p_{[2]}$ IS AT LEAST d^* . ($k = 4$)

The three values in each group are : (1) Normal approximation, (2) Straight line approximation, and (3) Smallest integer required.

d^*	P^*							
	0.50	0.60	0.75	0.80	0.85	0.90	0.95	0.99
0.05	69.85—	132.65+	282.27	357.52	456.82	599.53	848.30	1438.12
	70.02	132.99	282.98	358.42	457.96	601.03	850.42	1441.72
	71	134	283	359	458	601	850	1442
0.10	17.33	32.91	70.04	88.71	113.35—	148.76	210.48	356.83
	17.51	33.25—	70.74	89.61	114.49	150.26	212.61	360.43
	18	34	71	90	114	150	212	360
0.15	7.61	14.44	30.74	38.93	49.74	65.29	92.37	156.61
	7.78	14.78	31.44	39.82	50.88	66.78	94.49	160.19
	8	15	32	40	51	67	94	160
0.20	4.20	7.98	16.98	21.51	27.48	36.06	51.03	86.50+
	4.38	8.31	17.69	22.40	28.62	37.56	53.15+	90.12
	5	9	18	23	29	38	53	89
0.25	2.63	4.99	10.61	13.44	17.17	22.54	31.89	54.06
	2.80	5.32	11.32	14.34	18.32	24.04	34.02	57.67
	3	6	12	14	18	24	34	57
0.30	1.77	3.36	7.15+	9.06	11.58	15.19	21.50—	36.44
	1.95—	3.69	7.86	9.96	12.72	16.70	23.62	40.05—
	3	4	8	10	13	17	23	39
0.35	1.25+	2.38	5.07	6.42	8.20	10.76	15.23	25.82
	1.43	2.71	5.77	7.31	9.35—	12.27	17.36	29.42
	2	3	6	7	9	12	17	28
0.40	0.92	1.75—	3.71	4.70	6.01	7.89	11.16	18.92
	1.09	2.08	4.42	5.60	7.16	9.39	13.29	22.53
	2	3	5	6	7	9	13	21
0.45	0.69	1.31	2.79	3.53	4.51	5.92	8.37	14.19
	0.86	1.64	3.49	4.42	5.65+	7.42	10.51	17.80
	2	2	4	5	6	7	10	17
0.50	0.53	1.00	2.12	2.69	3.43	4.51	6.38	10.81
	0.70	1.33	2.83	3.58	4.58	6.01	8.50+	14.42
	2	2	3	4	5	6	8	13

TABLE IV — NUMBER OF UNITS REQUIRED PER PROCESS TO GUARANTEE A PROBABILITY OF P^* OF SELECTING THE BEST OF k BINOMIAL PROCESSES WHEN THE TRUE DIFFERENCE $p_{[1]} - p_{[2]}$ IS AT LEAST d^* . ($k = 10$)

The three values in each group are: (1) Normal approximation, (2) Straight line approximation, and (3) Smallest integer required.

d^*	P^*							
	0.50	0.60	0.75	0.80	0.85	0.90	0.95	0.99
0.05	216.96	312.51	511.15+	604.04	722.50−	887.54	1165.49	1798.01
	217.50+	313.29	512.43	605.55+	724.31	889.77	1168.41	1802.51
	218	314	513	606	725	890	1169	1803
0.10	53.83	77.54	126.83	149.87	179.27	220.22	289.18	446.12
	54.38	78.32	128.11	151.39	181.08	222.44	292.10	450.63
	55	79	128	151	181	222	291	449
0.15	23.62	34.03	55.66	65.77	78.67	96.64	126.90	195.77
	24.17	34.81	56.94	67.28	80.48	98.86	129.82	200.28
	25	35	57	67	80	98	129	198
0.20	13.05+	18.80	30.75−	36.33	43.46	53.39	70.10	108.15
	13.59	19.58	32.03	37.85−	45.27	55.61	73.03	112.66
	14	20	32	38	45	55	72	111
0.25	8.16	11.75−	19.22	22.71	27.16	33.37	43.82	67.59
	8.70	12.53	20.50−	24.22	28.97	35.59	46.74	72.10
	9	13	20	24	29	35	46	70
0.30	5.50−	7.92	12.95+	15.31	18.31	22.49	29.53	45.56
	6.04	8.70	14.23	16.82	20.12	24.72	32.46	50.07
	7	9	14	17	20	24	32	48
0.35	3.90	5.61	9.18	10.84	12.97	15.93	20.92	32.28
	4.44	6.39	10.46	12.36	14.78	18.16	23.85−	36.79
	5	7	11	13	15	18	23	35
0.40	2.85+	4.11	6.73	7.95−	9.51	11.68	15.34	23.66
	3.40	4.90	8.01	9.46	11.32	13.90	18.26	28.16
	4	5	8	10	11	13	17	26
0.45	2.14	3.08	5.05−	5.96	7.13	8.76	11.50+	17.75−
	2.69	3.87	6.33	7.48	8.94	10.98	14.42	22.25+
	3	4	6	8	9	11	14	20
0.50	1.63	2.35−	3.84	4.54	5.43	6.67	8.76	13.52
	2.18	3.13	5.12	6.06	7.24	8.90	11.68	18.03
	3	4	5	6	7	9	11	16

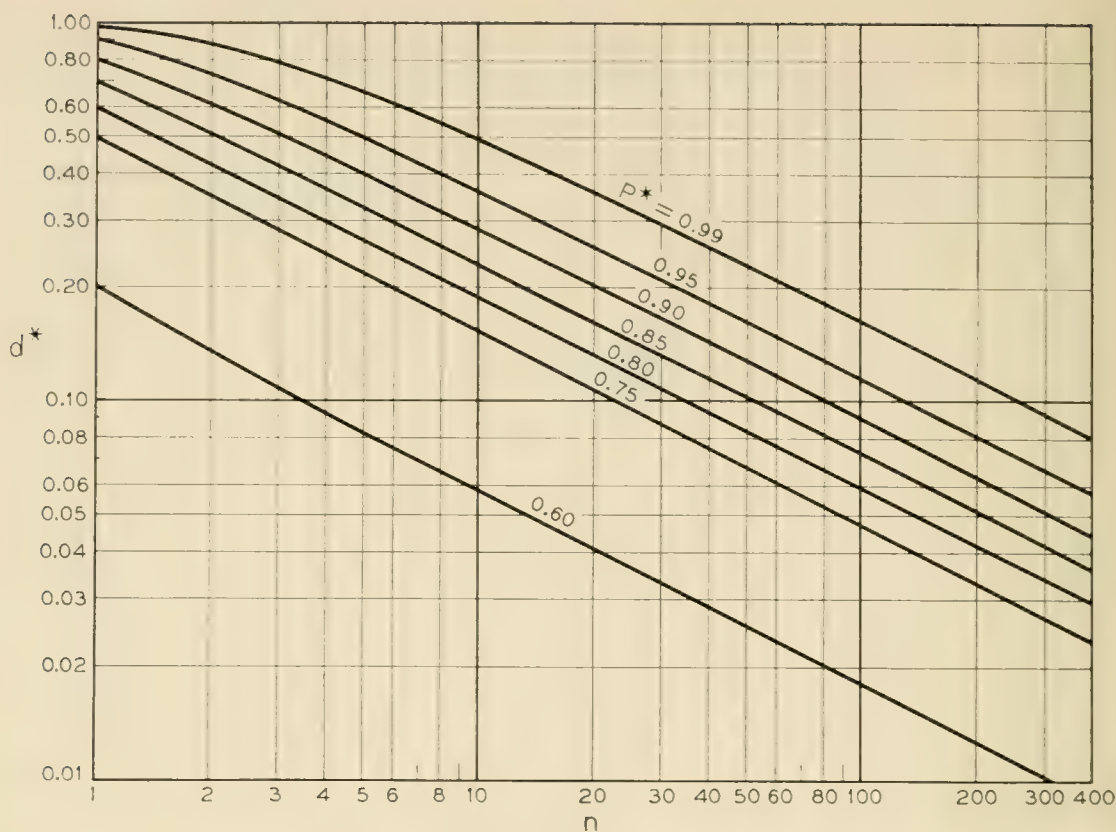


Fig. 2. — Number of units n required per process to guarantee a probability of P^* of selecting the better of two binomial processes when the true difference d is at least d^* .

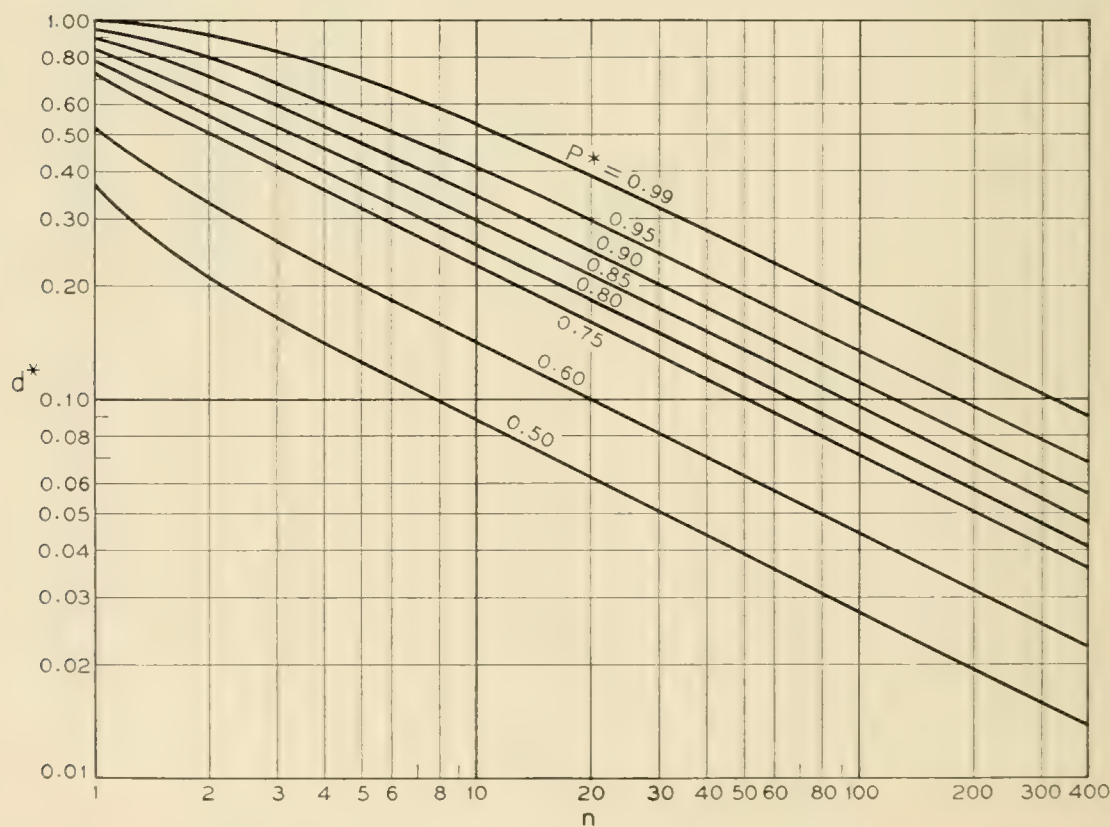


Fig. 3. — Number of units n required per process to guarantee a probability of P^* of selecting the best of three binomial processes when the true difference d is at least d^* .

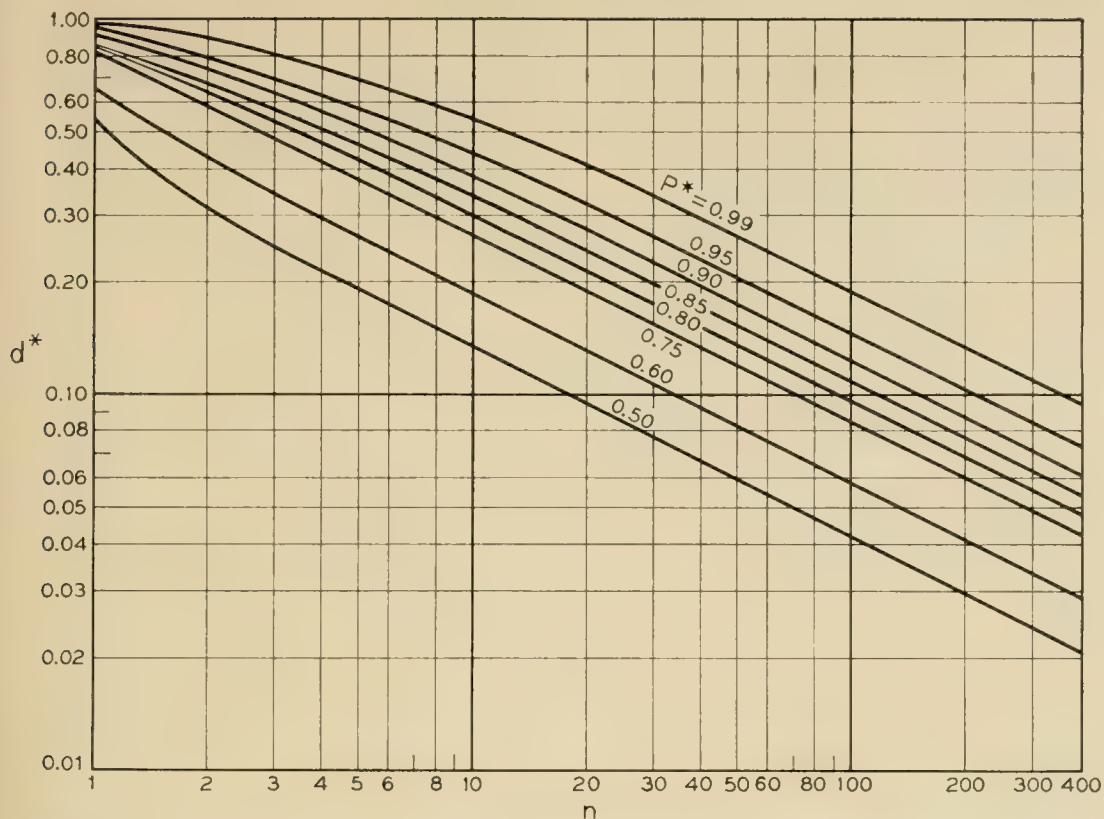


Fig. 4 — Number of units n required per process to guarantee a probability of P^* of selecting the best of four binomial processes when the true difference d is at least d^* .

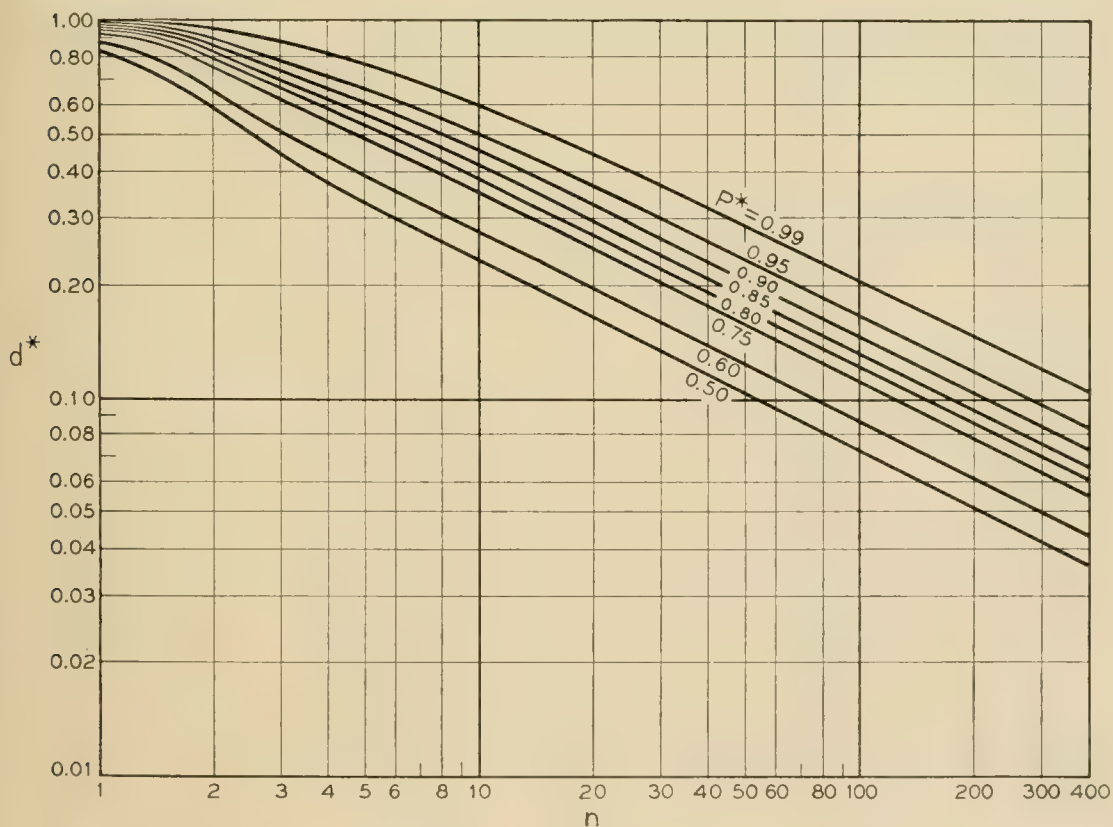


Fig. 5 — Number of units n required per process to guarantee a probability of P^* of selecting the best of ten binomial processes when the true difference d is at least d^* .

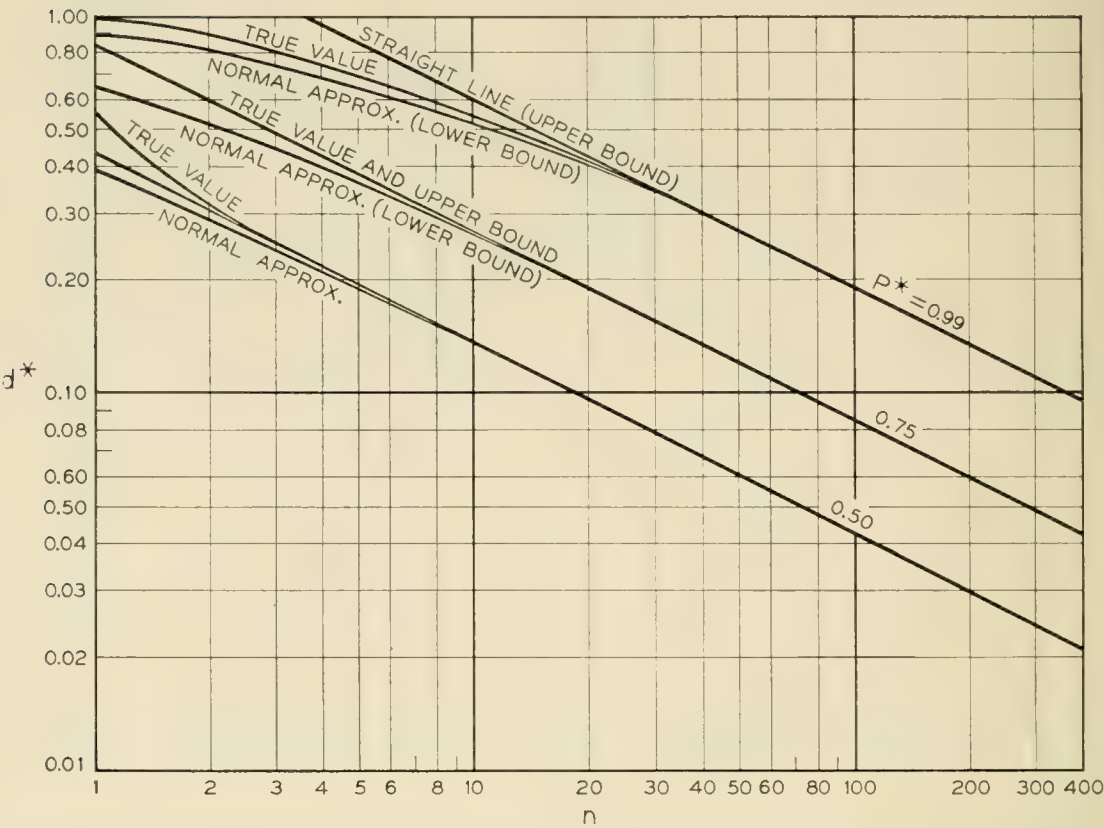


Fig. 6 — Bounds for the number of units n required per process to guarantee a probability of P^* of selecting the best of four binomial processes when the true difference d is at least d^* .

TABLE V — VALUES OF $B = \frac{1}{4}C^2$ TO BE USED WITH THE NORMAL APPROXIMATION (8) WHERE C IS OBTAINED FROM TABLE I OF R. E. BECHHOFFER'S PAPER¹

Prob. of Correct Selection	$k = 2$	$k = 3$	$k = 4$	$k = 10$
0.99	2.7060	3.2712	3.6043	4.5063
0.95	1.3528	1.8362	2.1261	2.9210
0.90	0.8212	1.2434	1.5026	2.2244
0.85	0.5371	0.9099	1.1449	1.7965+
0.80	0.3541	0.6826	0.8961	1.5139
0.75	0.2275—	0.5139	0.7074	1.2811
0.70	0.1375—	0.3832	0.5575—	1.0892
0.65	0.0742	0.2792	0.4347	0.9256
0.60	0.0321	0.1959	0.3325—	0.7832
0.55	0.0079	0.1294	0.2468	0.6569
0.50	0.0000	0.0774	0.1751	0.5438

order of magnitude of the error in our large sample calculations. For example, if $k = 4$, $d^* = 0.05$ and $P^* = 0.90$, then from Table V we find that $B = 1.5026$ and the two expressions in (8) yield 599.54 and 601.04. Hence, it would follow from (9) that n is 600 or 601 or 602. Based on an investigation of the behavior of these two approximations in the case of smaller P^* or larger d^* values, it is estimated that the true value of n is 601. Even if the correct value is 600 or 602 the error would be less than $\frac{1}{6}$ of 1 per cent. Fig. 6 illustrates these bounds on the P_{CS} for $k = 4$, $P^* = 0.50, 0.75$ and 0.99 . For $P^* \leq 0.60$ the straight line approximation is a closer lower bound than the normal approximation.

It is estimated that all integer entries in Tables I through IV have an error of at most 1 per cent and, in particular, that all entries under 100 are exact.

OTHER VALUES OF k

In addition to the tables and graphs for $k = 2, 3, 4$ and 10 there are also graphs (Figs. 7 through 14) on which interpolation can be carried out for $k = 5$ through 9 and on which extrapolation can be carried out for $k = 11$ through 100. By plotting n versus $\log k$ (or n versus k on semi-log paper) and drawing a straight (dashed) line through the values of n for $k = 4$ and $k = 10$ we obtain results which are remarkably good approximations for $k > 10$. The solid curve in these figures connects the true values obtained for $k = 2, 3, 4$ and 10.

For large values of k the theoretical justification for a straight line approximation is given in Appendix V. In order to check the accuracy of our procedure of drawing the straight line through the values of n computed for $k = 4$ and $k = 10$, we have chosen two points at $k = 101$ for an independent computation of the probability of a correct selection. For $P^* = 0.90$, $d^* = 0.10$ and $k = 101$ the dashed line in Fig. 12 gives n as approximately 400. To check this we computed the normal approximation to the probability P_{CS}^L of a correct selection for the least favorable configuration in the form

$$P_{CS}^L \cong \int_{-\infty}^{\infty} F^{100}(x + h)f(x) dx = \frac{1}{\sqrt{n}} \int_{-\infty}^{\infty} F^{100}(x\sqrt{2} + h)e^{-x^2} dx \quad (10)$$

where

$$h = \frac{2d^* \sqrt{n}}{\sqrt{1 - d^{*2}}} (= 4.02015 \text{ in this example}) \quad (11)$$

$f(x)$ is the normal density and $F(x)$ is its c.d.f. This was computed by a method suggested by Salzer, Zucker and Capuano³ and the result was

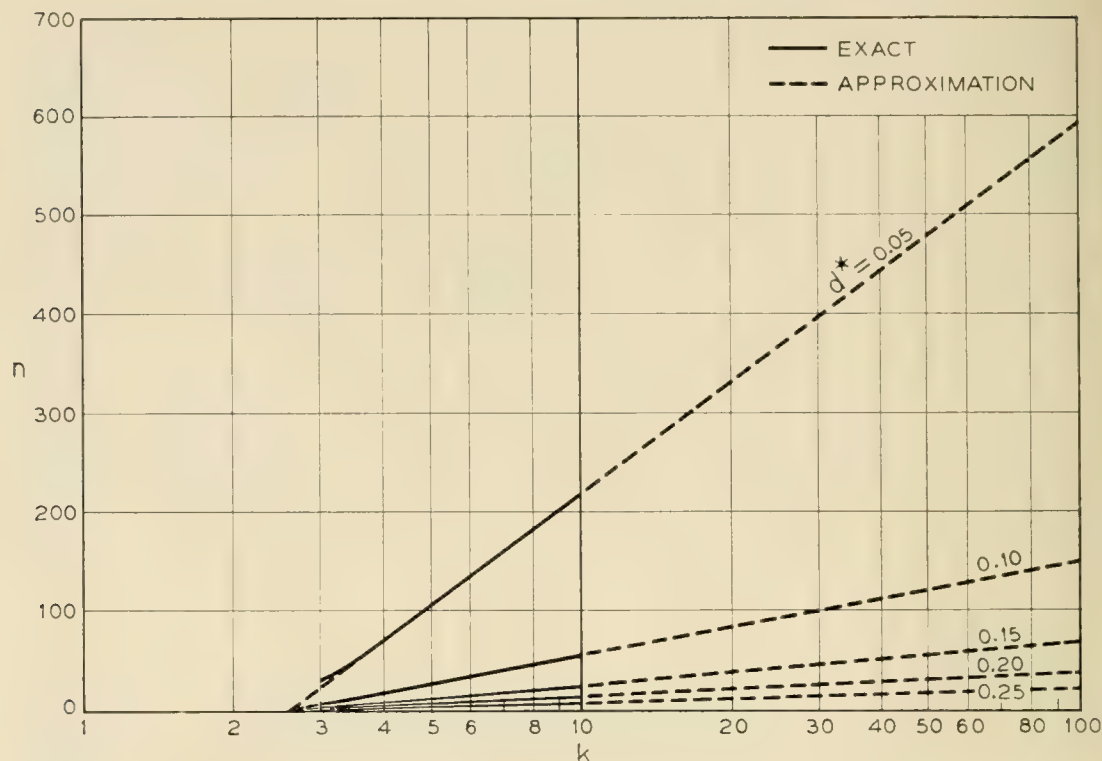


Fig. 7 — Number of units n required per process to guarantee a probability of $P^* = 0.50$ of selecting the best of k binomial processes when the true difference $p_{[1]} - p_{[2]}$ is at least d^* .

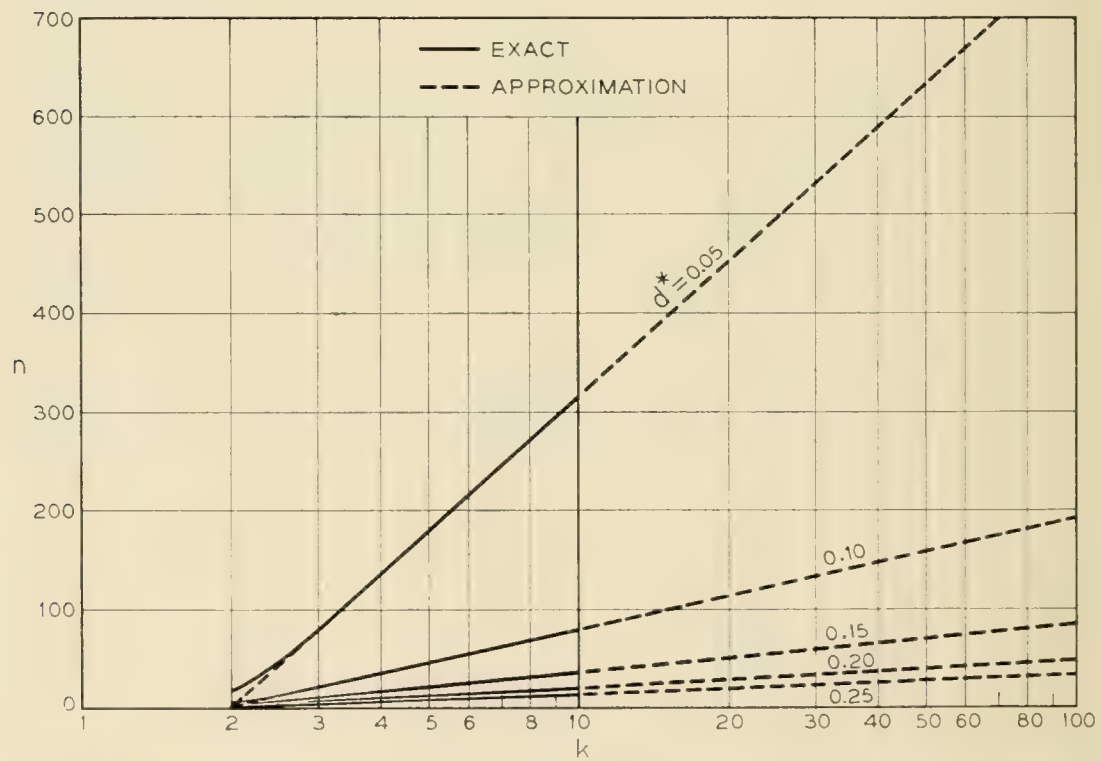


Fig. 8 — Number of units n required per process to guarantee a probability of $P^* = 0.60$ of selecting the best of k binomial processes when the true difference $p_{[1]} - p_{[2]}$ is at least d^* .

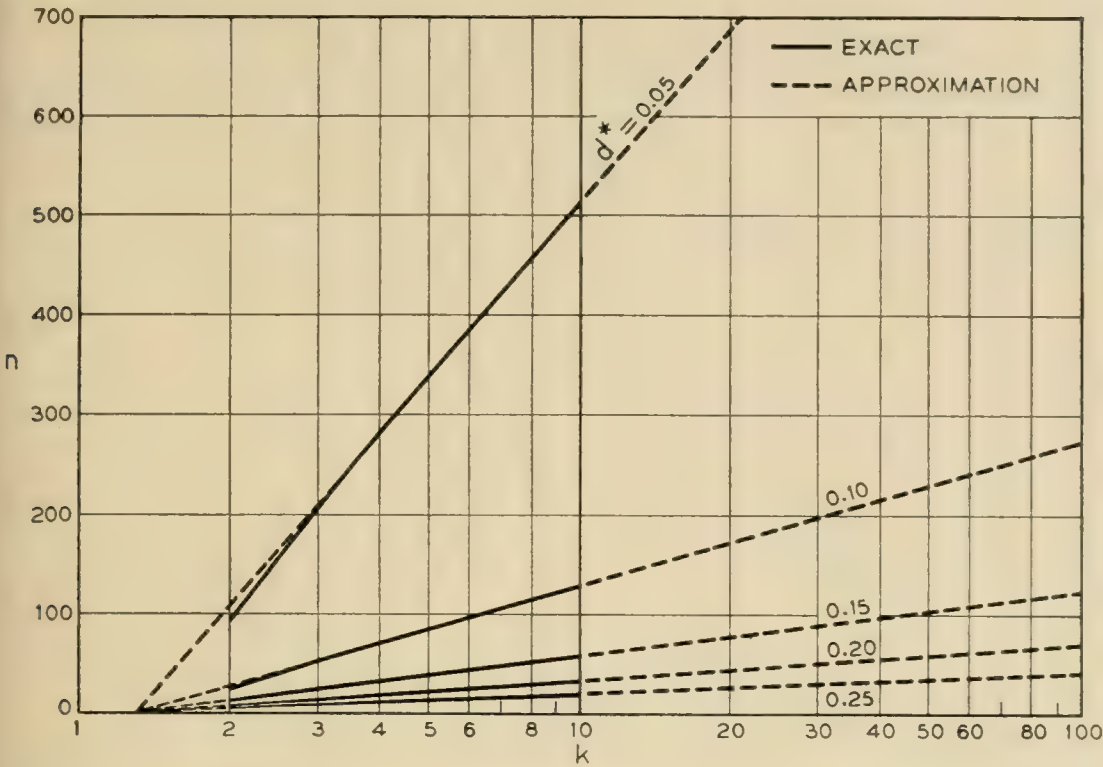


Fig. 9 — Number of units n required per process to guarantee a probability of $P^* = 0.75$ of selecting the best of k binomial processes when the true difference $p_{[1]} - p_{[2]}$ is at least d^* .

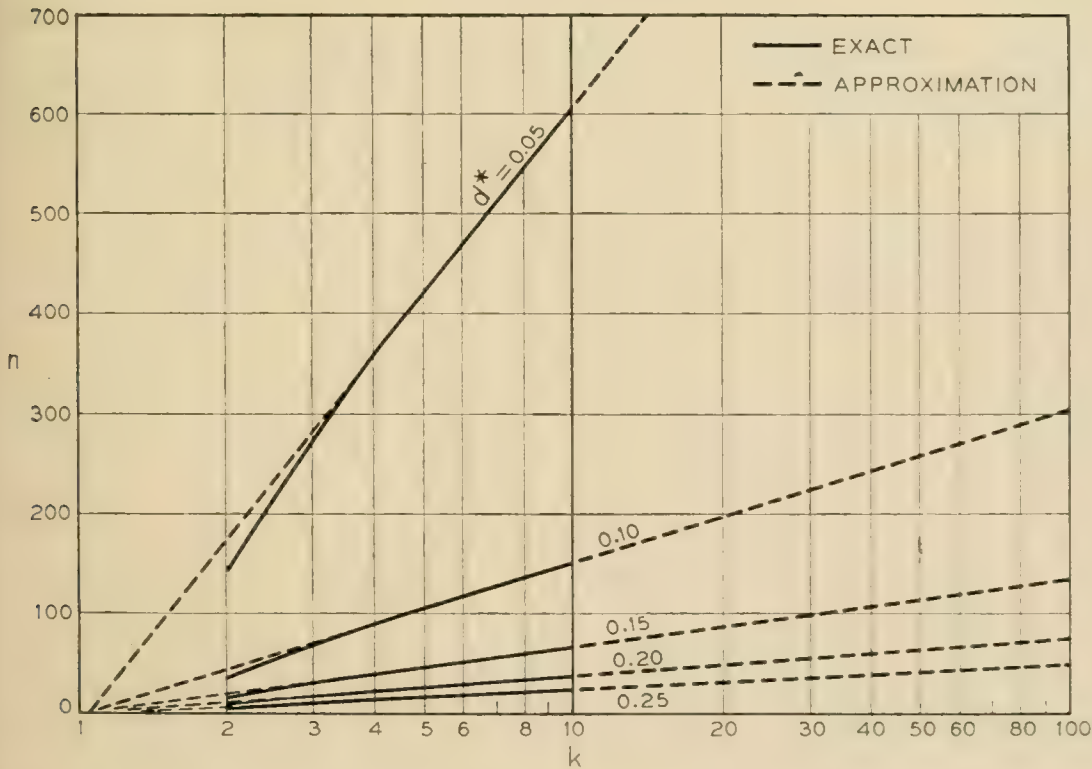


Fig. 10 — Number of units n required per process to guarantee a probability of $P^* = 0.80$ of selecting the best of k binomial processes when the true difference $p_{[1]} - p_{[2]}$ is at least d^* .

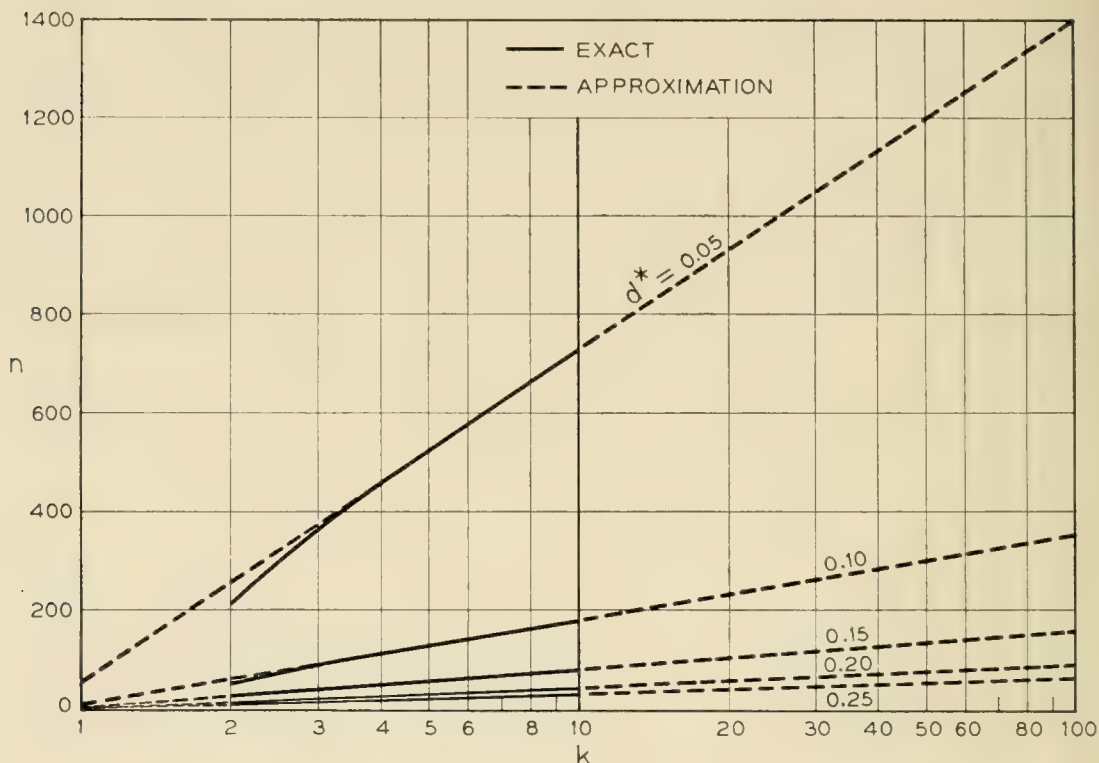


FIG. 11. — Number of units n required per process to guarantee a probability of $P^* = 0.85$ of selecting the best of k binomial processes when the true difference $p_{[1]} - p_{[2]}$ is at least d^* .

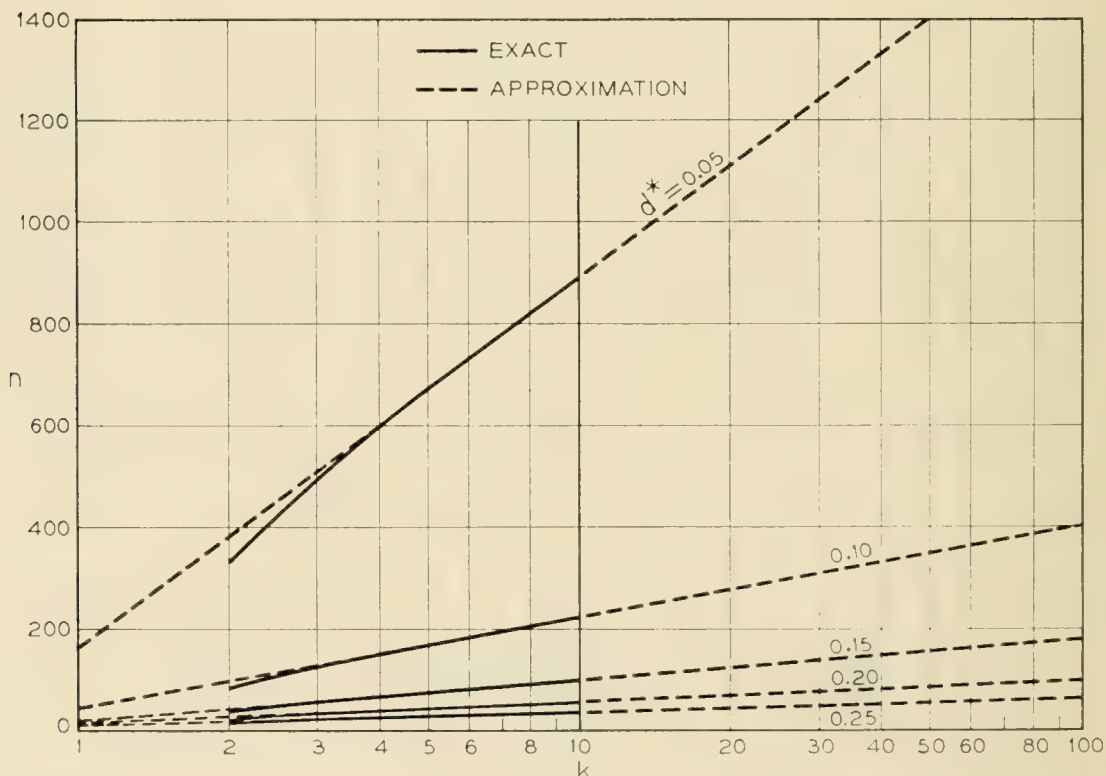


Fig. 12 — Number of units n required per process to guarantee a probability of $P^* = 0.90$ of selecting the best of k binomial processes when the true difference $p_{[1]} - p_{[2]}$ is at least d^* .

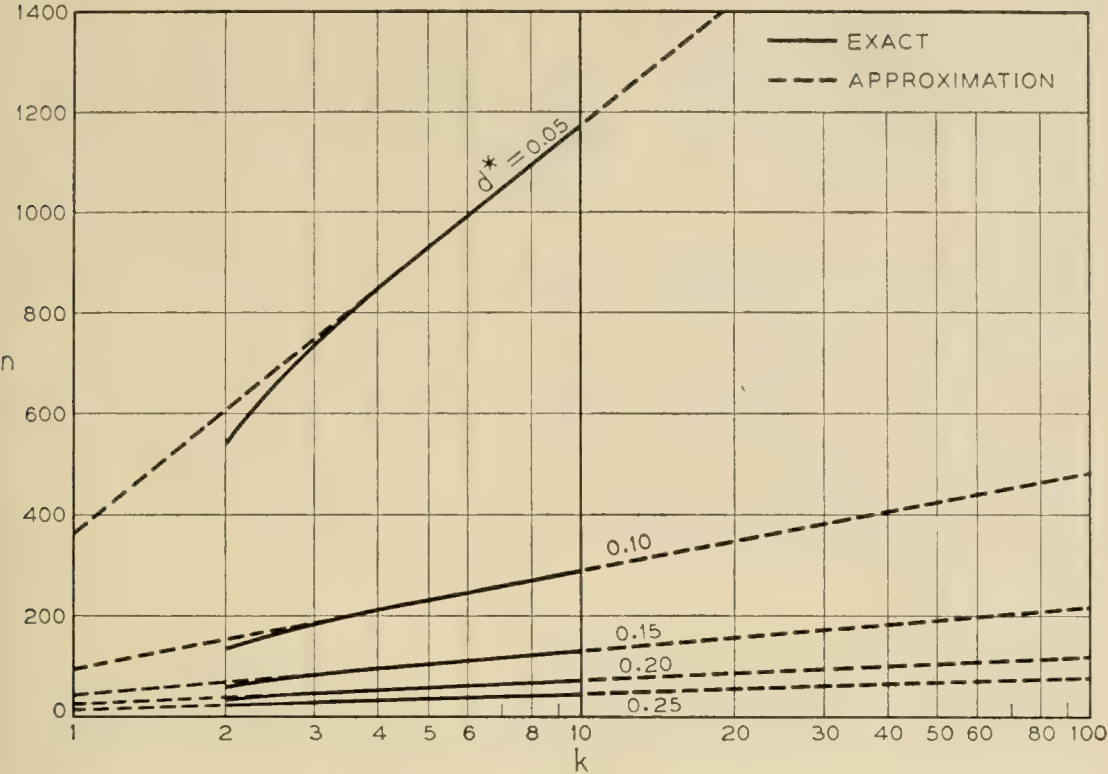


Fig. 13 — Number of units n required per process to guarantee a probability of $P^* = 0.95$ of selecting the best of k binomial processes when the true difference $p_{[1]} - p_{[2]}$ is at least d^* .

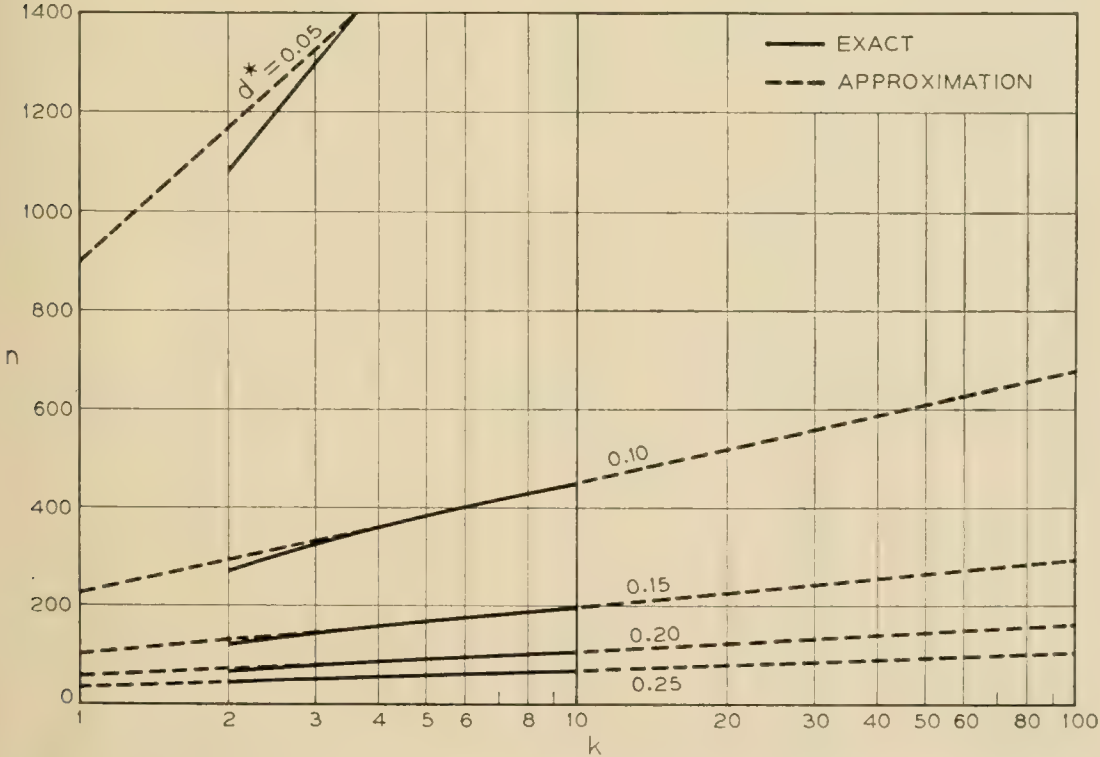


Fig. 14 — Number of units n required per process to guarantee a probability of $P^* = 0.99$ of selecting the best of k binomial processes when the true difference $p_{[1]} - p_{[2]}$ is at least d^* .

$P_{cs}^L \cong 0.9168$ as compared to the value 0.90 in Fig. 12. The expression (10) is derived in Appendix III. Another check was made at $P^* = 0.99$, $d^* = 0.20$ and $k = 101$. The value of n from Fig. 14 is 162. The value of the P_{cs}^L computed from (10) using Salzer, Zucker and Capuano³ is $0.9925+$.

Further calculation using (10) yielded the more accurate results 378 and 154 instead of 400 and 162, respectively, in the above illustrations. The error in both cases is less than 6 per cent; for smaller values of k the percentage error will, of course, be much less.

For interpolation the results are estimated to be within 1 per cent of the correct value. For example we estimate from Fig. 11 that the required value of n for $k = 5$, $P^* = 0.85$ and $d^* = 0.05$ is 523. This value was computed by the normal approximation and found to be 522.

TIED POPULATIONS

In computing the tables and graphs it was assumed that if two or more populations are tied for first place then one of *these* is selected by a chance device which assigns equal probability to each of them. The experimenter may want to select one of these contenders for first place by economic or other considerations. In most practical problems we may assume that such a selection is at random as far as the probability of a correct selection is concerned. Hence, it appears reasonable to use the tables in this paper without any corrections even when the rule for tied populations is altered in the manner described above.

It is interesting to note that in the yield problem the experimenter may settle the question of ties for first place by taking more observations until the tie is broken. However, in the life-testing problem he may not settle ties by letting the test run beyond time T since the best process for time T is not necessarily the best for a time greater than T .

In some applications when there are two or more populations tied for first place, the experimenter may prefer to recommend all these contenders for first place rather than select one of them by a chance device. In this case we shall agree to call the selection a correct one if the recommended set contains the best population (or, when $p_{[1]} = p_{[2]}$, if the recommended set contains at least one of the best populations). Exact tables for the procedure so altered have not been computed. However, if the value of n is large and this rule for tied populations is used, then the experimenter may reduce the tabled values by an amount equal to the largest integer contained in $1/d^*$. For example, using the above rule for tied populations for the case $k = 2$, $P^* = 0.99$, $d^* = 0.30$, the tabled value 29 can be reduced by 3 giving the result 26.

ALTERNATIVE SPECIFICATION

If the experimenter has some à priori knowledge about the processes, then he will prefer to specify the following *three* quantities in order to determine the number n of observations he is required to take from each process.

Specification: He specifies $p_{[1]}^*$ and $p_{[2]}^*$ ($0 \leq p_{[2]}^* \leq p_{[1]}^* \leq 1$) in the neighborhood of his estimate of the probabilities associated with his processes. He also specifies a probability P^* ($0 \leq P^* < 1$) that he would like to guarantee of making a correct selection whenever the true $p_{[1]} \geq p_{[1]}^*$ and the true $p_{[2]} \leq p_{[2]}^*$. (12)

TABLE VI — NUMBER OF UNITS REQUIRED PER PROCESS TO GUARANTEE A PROBABILITY OF P^* OF SELECTING THE BETTER OF TWO BINOMIAL PROCESSES WHEN THE TRUE $p_{[1]} \geq p_{[1]}^*$ AND THE TRUE $p_{[2]} \leq p_{[2]}^*$. (ALTERNATIVE SPECIFICATION, $k = 2$)

P^*	$p_{[1]}^* = 0.75$ $p_{[2]}^* = 0.60$	$p_{[1]}^* = 0.95$ $p_{[2]}^* = 0.80$	$p_{[1]}^* = 0.90$ $p_{[2]}^* = 0.80$	$p_{[1]}^* = 0.85$ $p_{[2]}^* = 0.80$	$p_{[1]}^* = 0.95$ $p_{[2]}^* = 0.90$
0.50	1	1	1	1	1
0.60	2	2	3	9	6
0.75	10	6	13	53	27
0.80	14	8	19	83	40
0.85	21	11	28	124	60
0.90	32	16	42	189	91
0.95	53	25	68	312	149
0.99	106	49	135	623	298

TABLE VII — NUMBER OF UNITS REQUIRED PER PROCESS TO GUARANTEE A PROBABILITY OF P^* OF SELECTING THE BEST OF FOUR BINOMIAL PROCESSES WHEN THE TRUE $p_{[1]} \geq p_{[1]}^*$ AND THE TRUE $p_{[2]} \leq p_{[2]}^*$. (ALTERNATIVE SPECIFICATION, $k = 4$)

P^*	$p_{[1]}^* = 0.75$ $p_{[2]}^* = 0.60$	$p_{[1]}^* = 0.95$ $p_{[2]}^* = 0.80$	$p_{[1]}^* = 0.90$ $p_{[2]}^* = 0.80$	$p_{[1]}^* = 0.85$ $p_{[2]}^* = 0.80$	$p_{[1]}^* = 0.95$ $p_{[2]}^* = 0.90$
0.50	7	4	10	42	21
0.60	14	8	18	79	39
0.75	28	14	37	168	80
0.80	35	18	46	211	101
0.85	45	22	59	268	128
0.90	59	28	77	350	171
0.95	83	39	107	493	239
0.99	139	65	182	831	399

Again we can rewrite the specification that the experimenter wants to satisfy in the simple form

$$P_{cs} \geq P^* \quad \text{for} \quad p_{[1]} \geq p_{[1]}^* \quad \text{and} \quad p_{[2]} \leq p_{[2]}^* \quad (13)$$

Tables VI and VII give the number of observations required per process for several selected triplets of specified constants ($p_{[1]}^*$, $p_{[2]}^*$, P^*) when $k = 2$ and $k = 4$. These results are also given in graphical form in Figs. 15 and 16.

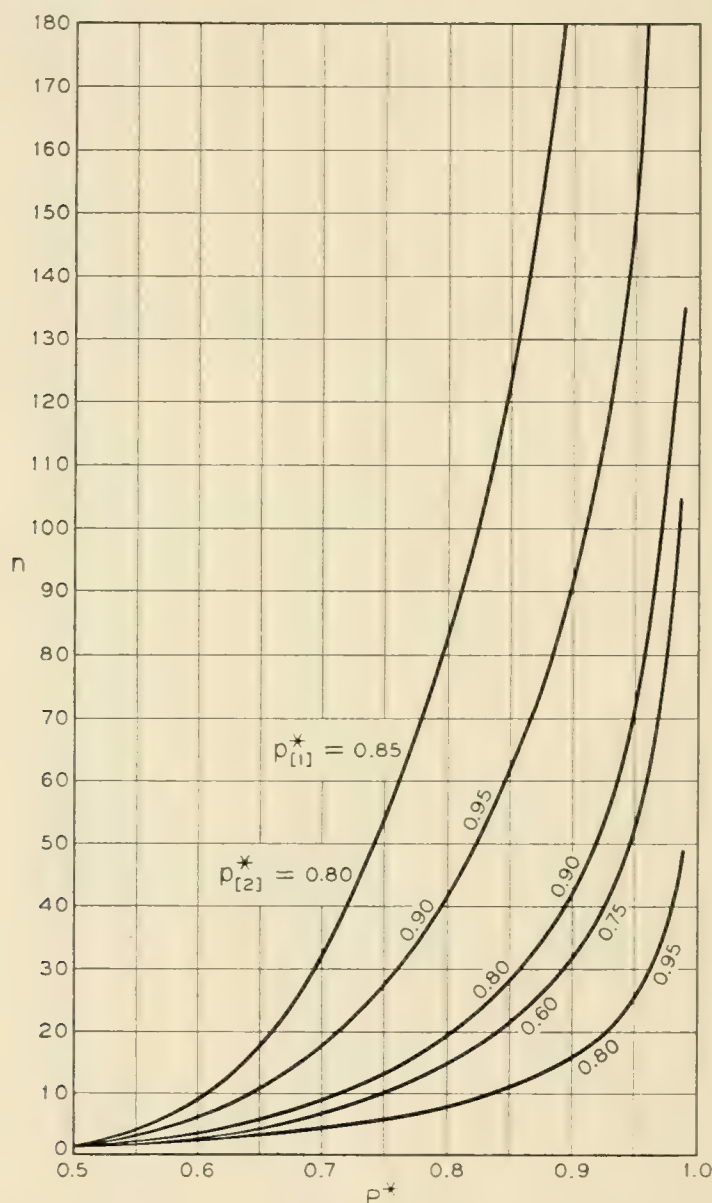


Fig. 15 — Number of units required per process to guarantee a probability of P^* of selecting the better of two binomial processes when the true $p_{[1]} \geq p_{[1]}^*$ and the true $p_{[2]} \leq p_{[2]}^*$.

For example, on the basis of past experience the experimenter may estimate that the probabilities associated with his $k = 4$ processes are all in the neighborhood of 0.60. This constitutes his à priori knowledge. He may then decide that he would like to make a correct selection with probability $P^* = 0.85$ when the best process has a yield of at least 75 per cent and all the others have a yield of at most 60 per cent. Entering column 1 of Table VII we find that $n = 45$ observations per process are required.

It is much more difficult to furnish tables for the alternative specifica-

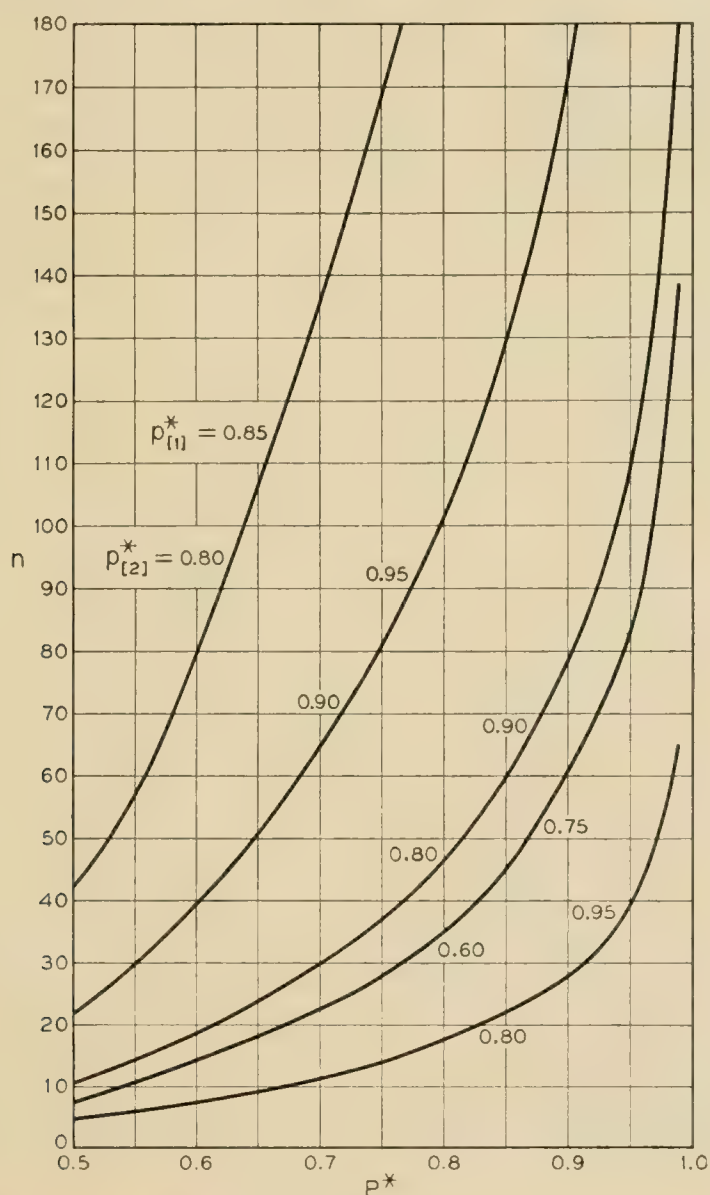


Fig. 16 — Number of units required per process to guarantee a probability of P^* of selecting the best of four binomial processes when the true $p_{[1]} \geq p_{[1]}^*$ and the true $p_{[2]} \leq p_{[2]}^*$.

tion since there is an extra parameter to vary and the appropriate tables for the normal approximation are not available.

In the computation of these probabilities the least favorable configuration

$$p_{[1]} = p_{[1]}^* \quad \text{and} \quad p_{[2]}^* = p_{[2]} = p_{[3]} = \cdots = p_{[k]} \quad (14)$$

was used. It follows from Appendix II that if the probability of a correct selection is at least P^* when (14) holds, then it will also be at least P^* when

$$p_{[1]} \geq p_{[1]}^* \quad \text{and} \quad p_{[2]}^* \geq p_{[2]} \geq p_{[3]} \geq \cdots \geq p_{[k]} \quad (15)$$

For small values of n , exact calculations were carried out. A typical exact calculation is shown in Appendix IV. The approximations used for large n are given in Appendix III.

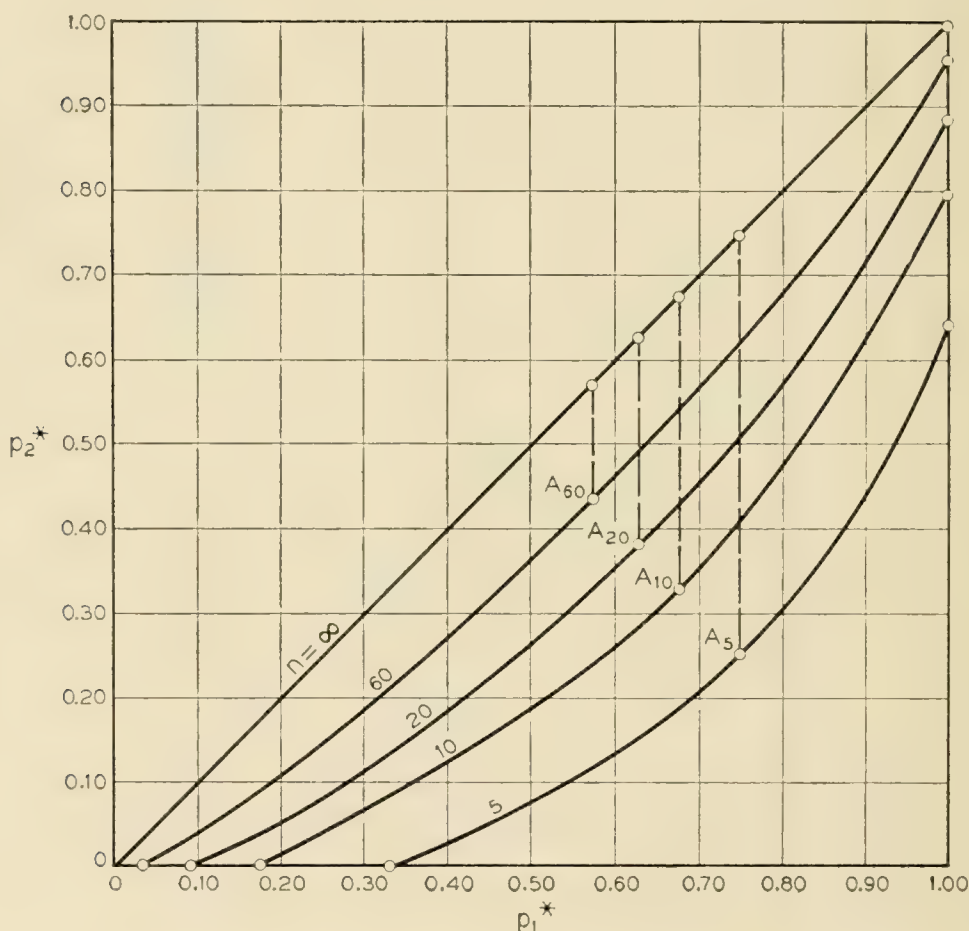


Fig. 17 — Illustration of the varying zones of indifference and the least favorable configuration for $k = 4$ and $P^* = 0.85$. (For $n \geq 5$ the longest vertical segment occurs at the point A_n where the abscissa and the ordinate are symmetrical about 0.5)

COMPARISON OF THE TWO SPECIFICATIONS

It should be pointed out that for a given k the same value of n would satisfy the specification for different specified triplets

$$P^*, p_{[1]}^*, p_{[2]}^*$$

For example with $k = 4$, $P^* = 0.85$ and n fixed we could vary $p_{[1]}^*$ in the alternative specification and compute for each $p_{[1]}^*$ the corresponding largest value of $p_{[2]}^*$ such that the specification $(P^*, p_{[1]}^*, p_{[2]}^*)$ is satisfied. This is shown in Fig. 17 for $n = 5, 10, 20$ and 60 . The vertical distance in Fig. 17 between the appropriate curve and the 45° line ($n = \infty$) is the length of the indifference zone $(p_{[1]}^*, p_{[2]}^*)$. The indifference zone widens in the center and narrows at both ends. In fact we find just as in the original specification that for n greater than (say) 4 the indifference zone is widest when $p_{[1]}^*$ and $p_{[2]}^*$ are symmetrical about 0.5. It is clear that the two specifications would coincide if we took d^* in the original specification and set $p_{[1]}^* = \frac{1}{2} (1 + d^*)$, $p_{[2]}^* = \frac{1}{2} (1 - d^*)$

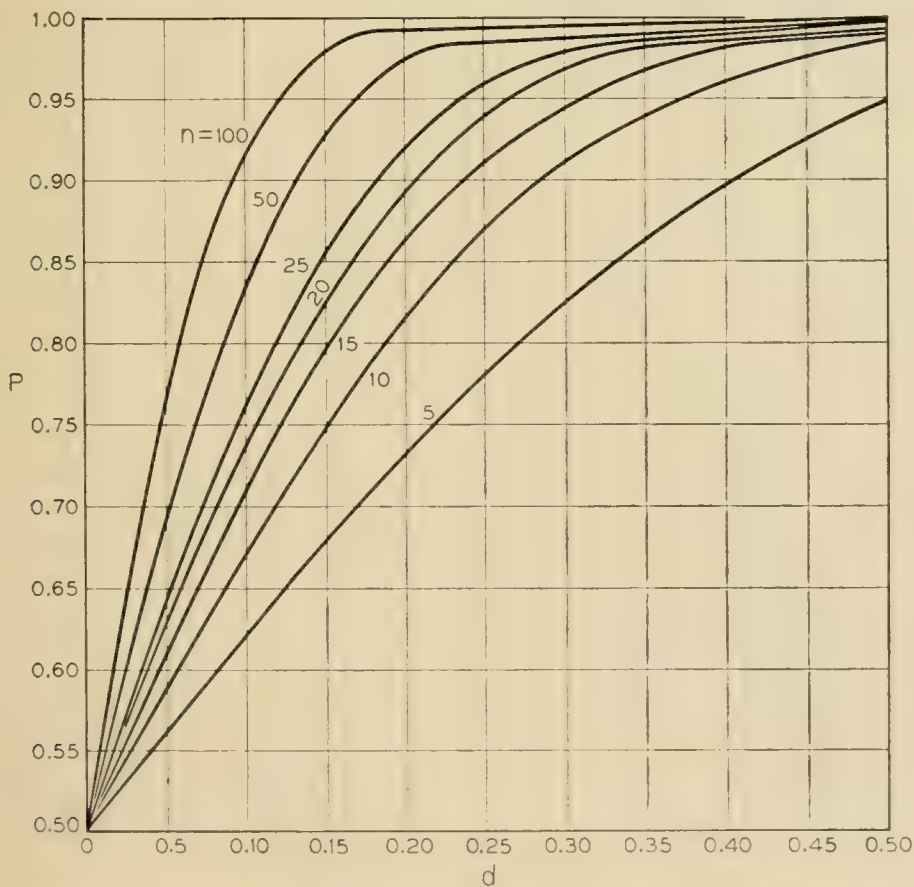


Fig. 18 — Probability of a correct selection as a function of the true difference $d = p_{[1]} - p_{[2]}$ under the least favorable configuration for $k = 2$ and selected values of n .

in the alternative specification. We shall be interested in comparing the alternative specification $(P^*, p_{[1]}^*, p_{[2]}^*)$ with the original specification with the same P^* and with d^* set equal to $p_{[1]}^* - p_{[2]}^*$. It is clear that the value of n required for the original specification will always be larger.

The original specification is simpler and is preferable to the alternative specification when little or nothing is known about the processes on test, but the price that has to be paid for ignorance is an increase in the

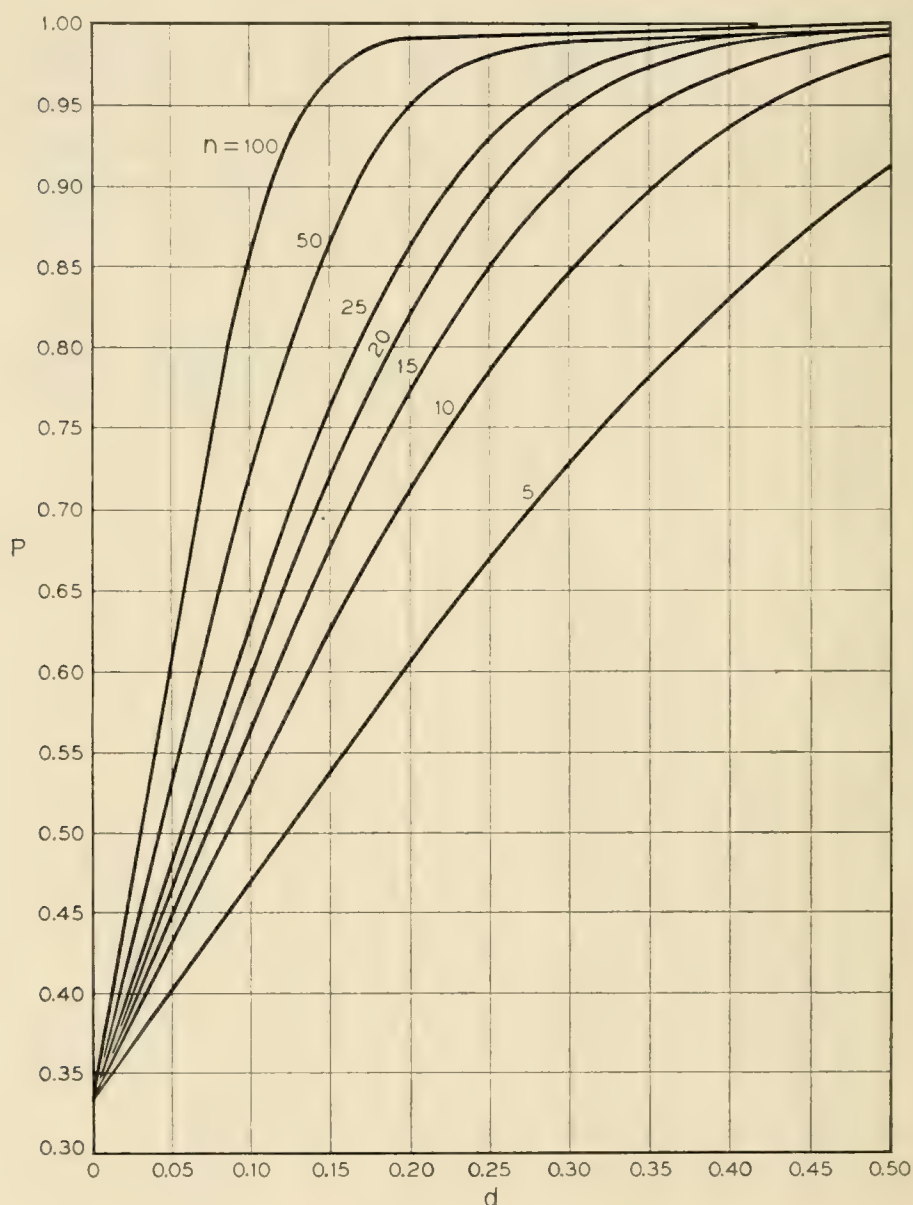


Fig. 19 — Probability of a correct selection as a function of the true difference $d = p_{[1]} - p_{[2]}$ under the least favorable configuration for $k = 3$ and selected values of n .

required number of observations. In the example of the preceding section the value of n required for the alternative specification is 45 as compared to 51 observations per process required to satisfy the original specification with the same P^* and with $d^* = p_{[1]}^* - p_{[2]}^*$. Here the saving is only moderate. The saving will be much larger if $p_{[1]}^*$ and

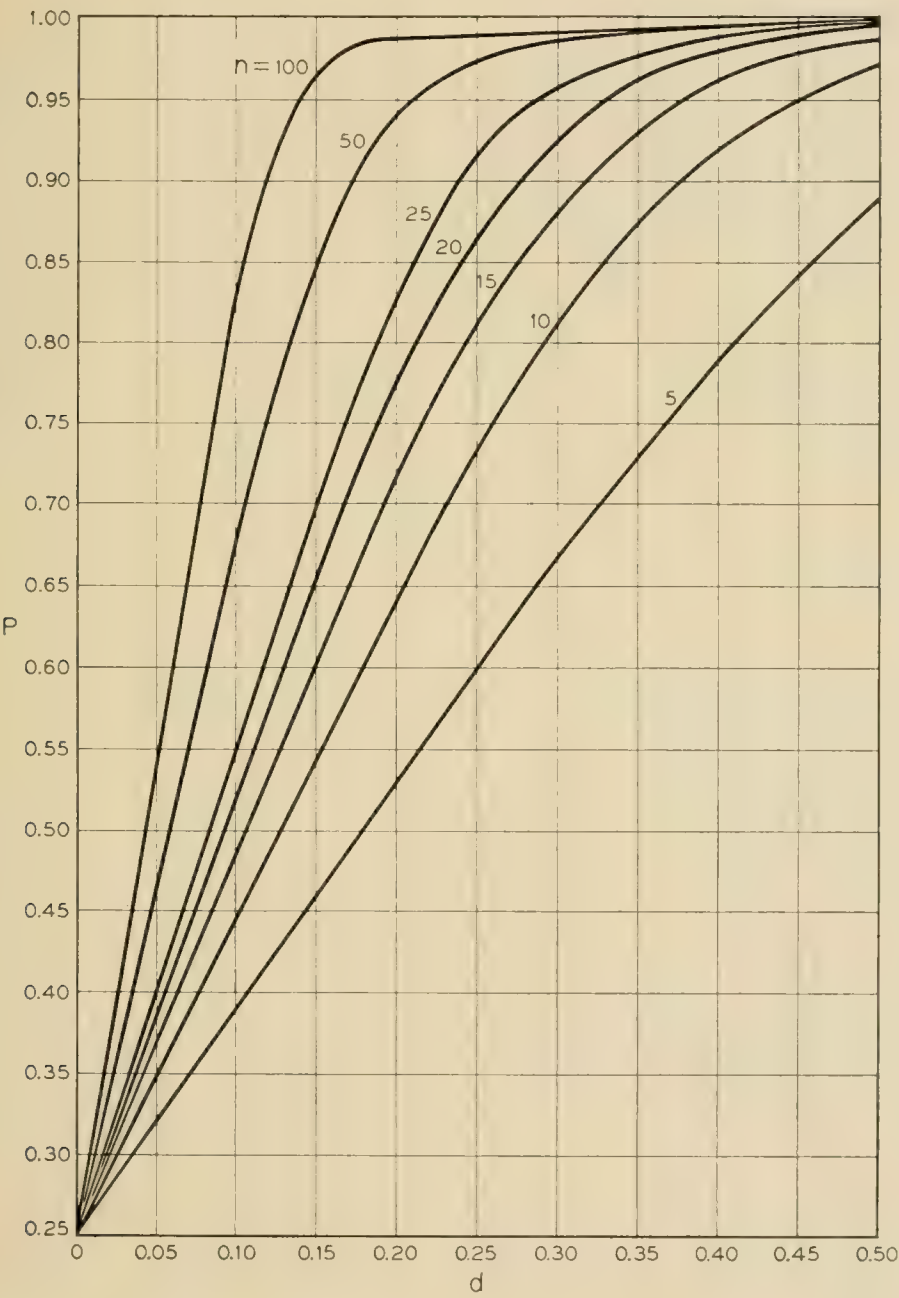


Fig. 20 — Probability of a correct selection as a function of the true difference $d = p_{[1]} - p_{[2]}$ under the least favorable configuration for $k = 4$ and selected values of n .

$p_{[2]}^*$ are further from 0.5 and d^* is small. For example, for $k = 4$, $P^* = 0.95$, $p_{[1]}^* = 0.95$ and $p_{[2]}^* = 0.90$ the value of n required for the alternative specification is 239 as compared to 850 observations per process required to satisfy the original specification with the same P^* and with $d^* = 0.05$. The alternative specification is justified on the basis of à priori or previous information about the approximate values of the p 's.

REVERSING THE TABLES

The experimenter may wish to use the tables of this paper in reverse. For example, if n is fixed and d^* is specified by the experimenter, then by using the appropriate table he can find the probability of a correct selection that is guaranteed for $d \geq d^*$; i.e., a greatest lower bound to the probability of a correct selection for $d \geq d^*$. This process of reversing the given values and the values to be computed can most easily be carried out on graphs. For example, the above problem of finding the guaranteed probability of a correct selection given d^* and n is most easily carried out on Figs. 18, 19, and 20.

APPENDIX I

MODIFICATION OF THE ORIGINAL SPECIFICATION

The same value of n will, of course, satisfy the specification for different pairs of specified values (d^* , P^*). From a purely mathematical point of view it is not necessary that d^* should be the *smallest* difference for which the experimenter desires to make a correct selection. For example, if $k = 3$ the experimenter could specify any one of the four pairs (0.10, 0.60), (0.25, 0.90), (0.30, 0.95) or (0.40, 0.99) and obtain the same result, namely $n = 20$. The experimenter may prefer to specify the *curve* or set of points corresponding to a fixed n . Several such curves are given in Figs. 18, 19, and 20 for $k = 2, 3$, and 4, respectively. The experimenter would decide in advance on some property of the curve that he considers desirable and from the appropriate figure he could find the curve with the smallest n -value that satisfies the desired property.

The main point of the above paragraph is to point out that the original specification in the body of the paper is one particular way, but not the only way, of stating a specification that will determine a value of n . The only criterion for a good way to state the specification is that the experimenter should be able to bring his best judgment (or best guesses) to bear on the quantities that have to be specified in advance.

APPENDIX II

MONOTONICITY PROPERTIES

We shall prove that for any fixed d ($0 \leq d \leq 1$) the probability P_{cs} of a correct selection is smaller for the configuration:

$$p_{[1]} - d = p_{[2]} = p_{[3]} = \cdots = p_{[k]} \tag{A1}$$

than for any configuration given by

$$p_{[1]} - d \geq p_{[2]} \geq p_{[3]} \geq \cdots \geq p_{[k]} \tag{A2}$$

where $p_{[1]}$ is considered fixed and the $p_{[i]}$ ($i \geq 2$) are variables. In other words, for fixed $p_{[1]}$ the probability P_{cs} is a strictly increasing function of each of the differences

$$p_{[1]} - p_{[i]} \tag{i \geq 2}$$

We shall need the following lemma.

Lemma 1: For any pair of integers x, n ($0 \leq x \leq n$) and any θ ($0 \leq \theta \leq 1$), not depending on p , the function

$$H(x; p, \theta) = \sum_{j=0}^{x-1} C_j^n p^j (1 - p)^{n-j} + \theta C_x^n p^x (1 - p)^{n-x} \tag{A3}$$

is a decreasing function of p over the unit interval ($0 \leq p \leq 1$). Moreover, it is strictly decreasing unless ($x = 0$ and $\theta = 0$) or ($x = n$ and $\theta = 1$).

Proof: Differentiating (A3) with respect to p gives after telescoping terms

$$(\theta - 1)x C_x^n p^{x-1} (1 - p)^{n-x} - \theta(n - x) C_x^n p^x (1 - p)^{n-x-1} \tag{A4}$$

which is negative for $0 < p < 1$ unless ($x = 0$ and $\theta = 0$) or ($x = n$ and $\theta = 1$). Since H is continuous in p at $p = 0$ and $p = 1$ the lemma follows.

Let $X_{(i)}$ denote the chance number of successes that arises from the binomial process associated with

$$p_{[i]} \tag{i = 1, 2, \cdots, n}$$

the value of the integer n is assumed to be fixed throughout this discussion and it will usually not be listed as an argument. The probability P_{cs} of a correct selection for any configuration with $p_{[1]} > p_{[2]}$ is given by the expression on the top of the next page.

$$P_{CS} = P\{X_{(i)} < X_{(1)} \text{ for } i \geq 2\} + \frac{1}{2} \sum_{\alpha=2}^k P\{X_{(\alpha)} = X_{(1)} \text{ and } X_{(i)} < X_{(1)} \text{ for } i \geq 2, i \neq \alpha\} + \dots$$

$$+ \frac{1}{k} P\{X_{(1)} = X_{(2)} = \dots = X_{(k)}\}$$
(A5)

It will be necessary to write the P_{CS} for any configuration with $p_{[1]} > p_{[2]}$ in another form which is more useful for the purpose at hand. Corresponding to any binomial chance variable X (which takes on integer values from 0 to n) we define a "Continuous Binomial" chance variable Y by letting Y be *uniformly* distributed in the interval $(j - \frac{1}{2}, j + \frac{1}{2})$ with the same total probability in this interval as the ordinary binomial assigns to the integer j , namely

$$C_j^n p^j (1-p)^{n-j} \quad (j = 0, 1, \dots, n)$$

We will now show that the probability P_{CS} of a correct selection is unaltered if we replace each of the k discrete binomials by its corresponding continuous binomial. Let $Y_{(i)}$ denote the continuous binomial (CB) chance variable associated with $p_{[i]}$ and let $y_{(i)}$ denote any value it can take on. Let $X_{(i)}$ denote the nearest integer to $Y_{(i)}$ and let $x_{(i)}$ denote the nearest integer to $y_{(i)}$ ($i = 1, 2, \dots, k$). Then $X_{(i)}$ is a discrete binomial (DB) with the same parameters ($p_{[i]}, n$). Let

$$g(x, p) = C_x^n p^x (1-p)^{n-x} \quad (x = 0, 1, \dots, n)$$

Then the *density* $g(y, p)$ of the continuous binomial (disregarding the half-integers) is given by $g(y, p) = g(x, p)$ where x is the nearest integer to y .

For two continuous binomials (i.e., $k = 2$) the probability P_{CS} of a correct selection for any configuration with $p_{[1]} > p_{[2]}$ is given by

$$P_{CS}(CB) = \int_{-1/2}^{n+1/2} P\{Y_{(2)} < y_{(1)}\} g(y_{(1)}; p_{[1]}) dy_{(1)} \quad (A6)$$

$$= \sum_{x_{(1)}=0}^n \int_{x_{(1)}-1/2}^{x_{(1)}+1/2} P\{Y_{(2)} < y_{(1)}\} g(y_{(1)}; p_{[1]}) dy_{(1)} \quad (A7)$$

Within any interval $(x_{(1)} - \frac{1}{2}, x_{(1)} + \frac{1}{2})$ we have

$$P\{Y_{(2)} < y_{(1)}\} = P\{X_{(2)} < x_{(1)}\} + P\{X_{(2)} = x_{(1)}\} P\{Y_{(2)} < y_{(1)} \mid X_{(2)} = x_{(1)}\}$$
(A8)

$$= P\{X_{(2)} < x_{(1)}\} + \frac{1}{2} P\{X_{(2)} = x_{(1)}\} \quad (A9)$$

which depends only on $x_{(1)}$. Hence from (A7)

$$\begin{aligned}
 P_{cs}(CB) &= \sum_{x_{(1)}=0}^n [P\{X_{(2)} < x_{(1)}\} + \tfrac{1}{2}P\{X_{(2)} = x_{(1)}\}]P\{X_{(1)} = x_{(1)}\} & (A10) \\
 &= P\{X_{(2)} < X_{(1)}\} + \tfrac{1}{2}P\{X_{(2)} = X_{(1)}\} & (A11) \\
 &= P_{cs}(DB) & (A12)
 \end{aligned}$$

The above is easily generalized to hold for any $k > 2$. The details of this generalization are omitted. For general k this equality holds not only for the important special case $p_{[1]} > p_{[2]}$ but also for the more general case (2) for any $t < k$. Since the latter result is not needed here, the proof is omitted.

If we let $G(y;p)$ denote the c.d.f. of the continuous binomial then lemma 1 can be restated in the following form.

Lemma 2: For any integer n and any y , the function $G(y;p)$ is a non-increasing function of p . In particular, for $-\frac{1}{2} < y < n + \frac{1}{2}$ it is a strictly decreasing function of p .

Proof: For any y , set $x = x(y)$ and $\theta = \theta(y)$ equal to the integer part and the fractional part of $(y + \frac{1}{2})$, respectively. Then for any y we have the identity in p

$$G(y;p) \equiv H(x;p, \theta) \qquad (0 \leq p \leq 1) \qquad (A13)$$

For any y_0 such that $-\frac{1}{2} < y_0 < n + \frac{1}{2}$ we have $0 \leq x(y_0) \leq n$ and $0 \leq \theta(y_0) \leq 1$. The inverse function $y(x,\theta) = x + \theta - \frac{1}{2}$ is a *single-valued* function of the pair (x,θ) ; the two particular pairs $(0,0)$ and $(n,1)$ correspond to the *unique* values $y = -\frac{1}{2}$ and $y = n + \frac{1}{2}$, respectively. Hence the pair $[x(y_0), \theta(y_0)]$ must be different from these two particular pairs above since it corresponds *only* to y_0 which is in the *interior* of the interval $(-\frac{1}{2}, n + \frac{1}{2})$. Lemma 2 then follows from lemma 1 and the fact that $G(y;p)$ is identically zero in p for $y \leq -\frac{1}{2}$ and identically one in p for $y \geq n + \frac{1}{2}$.

The probability P_{cs} of a correct selection for k discrete or k continuous binomials for any configuration with $p_{[1]} > p_{[2]}$ can now be written as

$$P_{cs} = \int_{-1/2}^{n+1/2} \left[\prod_{i=2}^k P\{Y_{(i)} < y_{(1)}\} \right] g(y_{(1)} ; p_{[1]}) dy_{(1)} \qquad (A14)$$

$$= \int_{-1/2}^{n+1/2} \left[\prod_{i=2}^k G(y;p_{[i]}) \right] g(y;p_{[1]}) dy. \qquad (A15)$$

Clearly if any one or more of the $p_{[i]}$ ($i \geq 2$) decreases, holding $p_{[1]}$ fixed, then it follows from lemma 2 that the right member of (A15) is

strictly increasing, i.e., for fixed $p_{[1]}$ the P_{CS} is a strictly increasing function of each of the differences $p_{[1]} - p_{[i]}$ ($i \geq 2$) as was to be shown.

It follows from the above result that in searching for a least favorable configuration among all those in which the experimenter wants his specification satisfied we may restrict our attention to those of the form (A1). Moreover, we may set d in (A1) equal to d^* since, for $d > d^*$ and fixed $p_{[1]}$, the difference $d - d^*$ may be added to each $p_{[i]}$ ($i \geq 2$) and the probability of a correct selection is increased. Then (A15) reduces to

$$P_{CS} = \int_{-1/2}^{n+1/2} G^{k-1}(y; p_{[1]} - d^*) g(y; p_{[1]}) dy \quad (\text{A16})$$

It was shown in the section on the least favorable configuration that there is a value $p_{[1]}^L$ of $p_{[1]}$ which when substituted in (A16) gives the minimum value P_{CS}^L of P_{CS} .

We can now prove the following result in which $p_{[1]}$ is not fixed. For any specified pair $p_{[2]}^* \leq p_{[1]}^*$ the probability P_{CS} of a correct selection is smaller for the configuration

$$p_{[1]} = p_{[1]}^*; \quad p_{[2]}^* = p_{[2]} = p_{[3]} = \cdots = p_{[k]} \quad (\text{A17})$$

than for any configuration given by

$$p_{[1]} \geq p_{[1]}^*; \quad p_{[2]}^* \geq p_{[2]} \geq p_{[3]} \geq \cdots \geq p_{[k]} \quad (\text{A18})$$

This is shown by considering two separate steps.

The first step is to increase $p_{[1]}$ holding all the other p 's fixed at $p_{[2]}^*$. For any arbitrary set of values of $p_{[i]}$ with $p_{[1]} > p_{[2]}$ the probability of a correct selection can be written as

$$P_{CS} = \sum_{j=2}^k \int_{-1/2}^{n+1/2} \left[\prod_{i=2, i \neq j}^k G(y, p_{[i]}) \right] [1 - G(y, p_{[1]})] g(y, p_{[j]}) dy \quad (\text{A19})$$

by adding the probabilities that

$$Y_{(1)} > Y_{(j)} > \min \{ Y_{(2)}, \cdots, Y_{(j-1)}, Y_{(j+1)}, \cdots, Y_{(k)} \}$$

for

$$j = 2, 3, \cdots, k$$

For

$$p_{[1]} > p_{[2]} = p_{[3]} = \cdots = p_{[k]} = p_{[2]}^*$$

this reduces to

$$P_{CS} = (k - 1) \int_{-1/2}^{n+1/2} [1 - G(y; p_{[1]})] G^{k-2}(y; p_{[2]}^*) g(y, p_{[2]}^*) dy \quad (\text{A20})$$

This result can also be obtained by starting with (A16) with

$$p_{[2]} = \cdots = p_{[k]} = p_{[1]} - d^* = p_{[2]}^*$$

and integrating by parts. It is clear from (A20) that for fixed $p_{[i]}$ ($i \geq 2$) the P_{cs} is an increasing function of $p_{[1]}$ and is indeed strictly increasing for $p_{[1]}$ in the unit interval.

The second step is to hold $p_{[1]}$ fixed and to decrease the values of $p_{[i]}$ ($i \geq 2$). This increases the probability of a correct selection by our previous result above. This proves the monotonicity property for the alternative specification.

APPENDIX III

LARGE SAMPLE THEORY — ORIGINAL SPECIFICATION

For $p_{[1]} > p_{[2]}$ the probability of a correct selection satisfies the inequalities

$$\begin{aligned} P\{X_{(1)} > X_{(i)} \ (i = 2, 3, \cdots, k)\} &< P_{cs} \\ &< P\{X_{(1)} \geq X_{(i)} \ (i = 2, 3, \cdots, k)\} \end{aligned} \tag{A21}$$

unless $p_{[1]} = 1$ and $p_{[2]} = 0$ in which case equality signs hold since the three quantities above are all unity. Letting $q_{[1]} = 1 - p_{[1]}$, we can write the left member of (A21) as

$$P\left\{Z_i > \frac{-d^* \sqrt{n}}{\sqrt{p_{[1]}q_{[1]} + (p_{[1]} - d^*)(q_{[1]} + d^*)}} \ (i = 1, 2, \cdots, k - 1)\right\} \tag{A22}$$

where

$$Z_i = \frac{X_{(1)} - X_{(i+1)} - nd^*}{\sqrt{n[p_{[1]}q_{[1]} + (p_{[1]} - d^*)(q_{[1]} + d^*)]} \ (i = 1, 2, \cdots, k - 1) \tag{A23}$$

For the configuration (A1) with $d = d^*$ the chance variables Z_i tend to normal chance variables $N(0,1)$ with zero mean and unit variance as $n \rightarrow \infty$. We have purposely omitted any continuity correction in (A22) in order to get a better approximation for the smaller values of n .

To derive the least favorable configuration for large n we can restrict our attention to those configurations given by (A1) with $d = d^*$. The quantity $p_{[1]}^L$, which minimizes (A22), is obtained by maximizing the expression in (A22)

$$Q(p) = p(1 - p) + (p - d^*)(1 - p + d^*) \tag{A24}$$

$$= -2p^2 + 2(1 + d^*)p - d^*(1 + d^*) \tag{A25}$$

The derivative of $Q(p)$ vanishes at

$$p_{[1]}^L = \frac{1}{2} (1 + d^*); \quad q_{[1]}^L = \frac{1}{2} (1 - d^*) \quad (\text{A26})$$

which gives the symmetric configuration. Clearly this value of p gives to $Q(p)$ its maximum value, $\frac{1}{2} (1 - d^{*2})$. This proves that the symmetrical configuration is least favorable in the limit as $n \rightarrow \infty$.

Under the configuration (A1) with $d = d^*$ and $n \rightarrow \infty$ the distribution of the chance variables Z_i ($i = 1, 2, \dots, k - 1$) approaches a joint multivariate normal distribution with zero means, unit variances and correlations given by

$$\rho(Z_i Z_j) = \frac{p_{[1]} q_{[1]}}{p_{[1]} q_{[1]} + (p_{[1]} - d^*)(q_{[1]} + d^*)} \quad (i \neq j) \quad (\text{A27})$$

which do not depend on n . For the symmetric configuration this reduces to the simple form

$$\rho(Z_i Z_j) = \frac{1}{2} \quad (i \neq j). \quad (\text{A28})$$

This is precisely the case which arises in [1] and consequently the tables in [1] can be used for our problem when (the answer) n is large. The constants $C = C(P^*, k)$ tabulated in [1] solve the equation

$$P \left\{ Z_i > -\frac{C}{\sqrt{2}} \quad (i = 1, 2, \dots, k - 1) \right\} = P^* \quad (\text{A29})$$

for standard normal chance variables Z_i satisfying (A28). If we equate $C/\sqrt{2}$ and the corresponding member of (A22), then we obtain for the symmetric configuration

$$\frac{C}{\sqrt{2}} \cong \frac{d^* \sqrt{n}}{\sqrt{\frac{1}{2} (1 - d^{*2})}} \quad (\text{A30})$$

or solving for n and letting $B = \frac{1}{4} C^2$ this yields the large sample normal approximation

$$n \cong \frac{B}{d^{*2}} (1 - d^{*2}) \quad (\text{A31})$$

Since d^* is usually small when n is large and since the solution in (A31) is usually somewhat smaller than the true value, then it is of interest to examine the simpler approximation

$$n \cong \frac{B}{d^{*2}} \quad (\text{A32})$$

which is greater than the result in (A31). This is called the straight line approximation since it plots as a straight line on log-log paper as shown

in Figs. 2 through 5. As $d^* \rightarrow 0$ both the normal approximation and the true value are asymptotically equivalent to the straight line approximation.

The normal approximation to the probability of a correct selection can also be written in another form similar to (A16) which is actually more useful for numerical calculations. The left member of (A21) can be written as

$$\sum_{w_1} \left[\prod_{i=2}^{k-1} P \left\{ W_i < \frac{W_1 \sqrt{p_{[1]}q_{[1]}} + (p_{[1]} - p_{[i]}) \sqrt{n}}{\sqrt{p_{[i]}q_{[i]}}} \right\} \right] P \{ W_1 = w_1 \} \tag{A33}$$

where

$$W_i = \frac{X_{(i)} - np_{[i]}}{\sqrt{np_{[i]}q_{[i]}}} \qquad (i = 1, 2, \dots, k) \tag{A34}$$

and w_1 is the same function of $x_{(1)}$ as W_1 is of $X_{(1)}$. The outside summation in (A33) is over the values taken on by w_1 as $x_{(1)}$ runs from 0 to n . As $n \rightarrow \infty$ the expression in (A33) approaches

$$P_{cs} \cong \int_{-\infty}^{\infty} \left[\prod_{i=2}^{k-1} F \left(\frac{w \sqrt{p_{[1]}q_{[1]}} + (p_{[1]} - p_{[i]}) \sqrt{n}}{\sqrt{p_{[i]}q_{[i]}}} \right) \right] f(w) dw \tag{A35}$$

where $f(t)$ is the standard normal density and $F(t)$ is the standard normal c.d.f. For the symmetric configuration, which is least favorable for large n , (A35) reduces to

$$P_{cs}^L \cong \int_{-\infty}^{\infty} F^{k-1} \left(w + \frac{2d^* \sqrt{n}}{\sqrt{1 - d^{*2}}} \right) f(w) dw \tag{A36}$$

A straightforward integration by parts gives the alternative form

$$P_{cs}^L \cong (k - 1) \int_{-\infty}^{\infty} \left[1 - F \left(w - \frac{2d^* \sqrt{n}}{\sqrt{1 - d^{*2}}} \right) \right] F^{k-2}(w) f(w) dw \tag{A37}$$

which corresponds to (A20).

A simple method for computing such integrals based on Hermite polynomials is described by Salzer, Zucker, and Capuano.³

LARGE SAMPLE THEORY — ALTERNATIVE SPECIFICATION

The expression corresponding to (A22) for the alternative specification is

$$P \left\{ Z_i > \frac{-(p_{[1]}^* - p_{[2]}^*) \sqrt{n}}{\sqrt{p_{[1]}^*q_{[1]}^* + p_{[2]}^*q_{[2]}^*}} \qquad (i = 1, 2, \dots, k - 1) \right\} \tag{A38}$$

which is already written for the least favorable configuration. The tables² are not immediately applicable since the correlations

$$\rho(Z_i Z_j) = \frac{p_{[1]}^* q_{[1]}^*}{p_{[1]}^* q_{[1]}^* + p_{[2]}^* q_{[2]}^*} \quad (i \neq j) \quad (\text{A39})$$

are not, in general, equal to $\frac{1}{2}$. In the cases treated in Tables VI and VII, $p_{[2]}^* \geq 0.5$ and hence $p_{[1]}^* q_{[1]}^* < p_{[2]}^* q_{[2]}^*$ so that the correlations (A39) are all less than $\frac{1}{2}$. It was found that linear interpolation on the required value of n between the results for $\rho = 0$ and $\rho = \frac{1}{2}$ gives moderately good results when n is large. The result for $\rho = \frac{1}{2}$ is given by

$$n \cong \lambda \frac{(p_{[1]}^* q_{[1]}^* + p_{[2]}^* q_{[2]}^*)}{(p_{[1]}^* - p_{[2]}^*)^2} \quad (\text{A40})$$

with $\lambda = 2B$ where B is given in Table V. The result for $\rho = 0$ is given by (A40) with $\lambda = \lambda_0^2$ where λ_0 is the solution of the equation

$$P\{Z > -\lambda_0\} = P^{*1/(k-1)} \quad (\text{A41})$$

which can easily be found from univariate normal probability tables. An explicit expression for the result of this linear interpolation is

$$P_{cs}^L \cong \frac{p_{[1]}^* q_{[1]}^* (4B - \lambda_0^2) + p_{[2]}^* q_{[2]}^* \lambda_0^2}{(p_{[1]}^* - p_{[2]}^*)^2} \quad (\text{for } p_{[2]}^* \geq 0.5) \quad (\text{A42})$$

The expressions for the probability of a correct selection for the alternative specification corresponding to (A36) and (A37) are

$$P_{cs}^L \cong \int_{-\infty}^{\infty} F^{k-1}(aw + b) f(w) dw \quad (\text{A43})$$

$$= a(k-1) \int_{-\infty}^{\infty} \left[1 - F\left(\frac{w-b}{a}\right) \right] F^{k-2}(w) f(w) dw \quad (\text{A44})$$

where

$$a = \sqrt{\frac{p_{[1]}^* q_{[1]}^*}{p_{[2]}^* q_{[2]}^*}} > 0 \quad \text{and} \quad b = \frac{(p_{[1]}^* - p_{[2]}^*) \sqrt{n}}{\sqrt{p_{[2]}^* q_{[2]}^*}} \geq 0 \quad (\text{A45})$$

These expressions can also be evaluated by the method described by Salzer, Zucker and Capuano.³

APPENDIX IV

TYPICAL EXACT CALCULATION

A. Original Specification

The exact expression (A5) for the probability of a correct selection for any configuration simplifies if the configuration is least favorable. For any pair of integers (j, n) we define

$$b_{1j} = P\{X_{(1)} = j\} = C_j^n (p_{[1]}^L)^j (q_{[1]}^L)^{n-j} \quad (0 \leq j \leq n) \quad (\text{A46})$$

$$b_{2j} = P\{X_{(2)} = j\} = C_j^n (p_{[1]}^L - d^*)^j (q_{[1]}^L + d^*)^{n-j} \quad (0 \leq j \leq n) \quad (\text{A47})$$

$$B_{2j} = P\{X_{(2)} \leq j\} \quad (\text{A48})$$

Then the exact probability P_{CS}^L of a correct selection for the least favorable configuration can be written as

$$P_{CS}^L = \sum_{j=0}^n b_{1j} \sum_{i=0}^{k-1} \frac{C_i^{k-1}}{1+i} b_{2j}^i B_{2,j-1}^{k-1-i} \quad (\text{A49})$$

where $B_{2,-1}$ is defined to be zero. Here, for each value of $X_{(1)}$, the letter i denotes the number of processes that tie with $X_{(1)}$ for first place and for any given value of i the conditional probability of a correct selection is $1/(1+i)$. Taking $k=4$ as a typical case, we can write (A49) more explicitly as

$$\begin{aligned} P_{CS}^L = & \sum_{j=1}^n b_{1j} B_{2,j-1}^3 + \frac{3}{2} \sum_{j=1}^n b_{1j} b_{2j} B_{2,j-1}^2 + \sum_{j=1}^n b_{1j} b_{2j}^2 B_{2,j-1} \\ & + \frac{1}{4} \sum_{j=0}^n b_{1j} b_{2j}^3 \end{aligned} \quad (\text{A50})$$

If $n \geq 10$ then we may use the symmetric configuration, i.e., we may set $p_{[1]}^L = \frac{1}{2}(1+d^*)$, in computing from (A49) or (A50).

B. Alternative Specification

The probability P_{LS}^C of a correct selection for the alternative specification is the same as in (A49) and (A50) except that we now define

$$b_{ij} = P\{X_{(i)} = j\} = C_j^n (p_{[i]}^*)^j (q_{[i]}^*)^{n-j} \quad (i = 1, 2) \quad (\text{A51})$$

$$B_{2j} = P\{X_{(2)} \leq j\} \quad (\text{A52})$$

A typical exact calculation for $k=4$, using (A50), (A51) and (A52) with

$$p_{[1]} = p_{[1]}^* = 0.75$$

and

$$p_{[2]} = p_{[3]} = p_{[4]} = p_{[2]}^*$$

is given in Table AI. Exact values for the individual and cumulative binomial probabilities were obtained from References 4, 5 and 6.

TABLE AI — CALCULATION OF THE P_{cs}^L							
$p_{[1]} = p_{[1]}^* = 0.75; \quad p_{[2]} = p_{[3]} = p_{[4]} = p_{[2]}^* = 0.60; \quad k = 4$							
j	$b_{1,j}$	$b_{2,j}$	$B_{2,j-1}$	$b_{2,j}^2$	$b_{2,j}^3$	$B_{2,j-1}^2$	$B_{2,j-1}^3$
0	0.00000	0.00010	—	0.00000	0.00000	—	—
1	0.00003	0.00157	0.00010	0.00000	0.00000	0.00000	0.00000
2	0.00039	0.01062	0.00168	0.00011	0.00000	0.00000	0.00000
3	0.00309	0.04247	0.01229	0.00180	0.00008	0.00015	0.00000
4	0.01622	0.11148	0.05476	0.01243	0.00139	0.00300	0.00016
5	0.05840	0.20066	0.16624	0.04026	0.00808	0.02764	0.00459
6	0.14600	0.25082	0.36690	0.06291	0.01578	0.13462	0.04939
7	0.25028	0.21499	0.61772	0.04622	0.00994	0.38158	0.23571
8	0.28157	0.12093	0.83271	0.01462	0.00177	0.69341	0.57741
9	0.18771	0.04031	0.95364	0.00162	0.00007	0.90943	0.86727
10	0.05631	0.00605	0.99395	0.00004	0.00000	0.98794	0.98196
Check totals....	1.00000	1.00000					

$$\sum_{j=1}^{10} b_{1,j} B_{2,j-1}^3 = 0.44715$$

$$\frac{3}{2} \sum_{j=1}^{10} b_{1,j} b_{2,j} B_{2,j-1}^2 = 0.08493$$

$$\sum_{j=1}^{10} b_{1,j} b_{2,j}^2 B_{2,j-1} = 0.01464$$

$$\frac{1}{4} \sum_{j=0}^{10} b_{1,j} b_{2,j}^3 = 0.00145$$

$$\text{Total} = \overline{0.54817} = P_{cs}^L$$

APPENDIX V

In this appendix it will be shown that for large values of k the value of n required to meet any fixed specification (d^* , P^*) is approximately equal to some constant multiple of $(\ln k)$.

Let $n = n(k)$ denote the unique positive decimal solution of the equation

$$\int_{-\infty}^{\infty} F^{k-1}(w + b \sqrt{n})f(w) dw = P^* \tag{A53}$$

where $f(w)$ and $F(w)$ are defined above, P^* and b are known constants with $1/k < P^* < 1$ and $b > 0$ and the argument k is a positive integer. Let ε be a (small) fixed number such that $0 < \varepsilon < \text{Min}(P^*, 1 - P^*)$. Then $\varepsilon < P^* - 1/k$ for sufficiently large k . Let $A = A(\varepsilon)$ be defined by

$$\int_{-A}^A f(w) dw = 1 - \varepsilon \quad (\text{A54})$$

so that

$$0 < \int_{|w| > A} F^{k-1}(w + b\sqrt{n})f(w) dw \leq \varepsilon \quad (\text{A55})$$

for any integer $k \geq 1$, any $n > 0$ and any $b > 0$. Let n' and n'' be the unique positive decimal solutions, respectively, of the equations

$$\int_{-A}^A F^{k-1}(w + b\sqrt{n'})f(w) dw = P^* - \varepsilon \quad (\text{A56})$$

$$\int_{-A}^A F^{k-1}(w + b\sqrt{n''})f(w) dw = P^* \quad (\text{A57})$$

where P^* , b and k are the same as in (A53). It follows from (A55), (A56) and (A57) that for any integer $k \geq 1$

$$n' \leq n \leq n'' \quad (\text{A58})$$

From (A54) and (A57) we have

$$\int_{-A}^A F^{k-1}(w + b\sqrt{n''})f_A(w) dw = \frac{P^*}{1 - \varepsilon} \quad (\text{A59})$$

where $f_A(w)$ is the density of the normal distribution, truncated at A and $-A$. The right hand member of (A59) is positive and less than unity since $\varepsilon < 1 - P^*$. Hence there exists a w_A with $|w_A| \leq A$ such that

$$\int_{-A}^A F^{k-1}(w + b\sqrt{n''})f_A(w) dw = F^{k-1}(w_A + b\sqrt{n''}) \quad (\text{A60})$$

Since w_A is bounded and n'' is large for large k we can use the well-known approximation

$$\begin{aligned} F^{k-1}(w_A + b\sqrt{n''}) &\cong \left[1 - \frac{\exp [-(w_A + b\sqrt{n''})^2/2]}{\sqrt{2\pi}(w_A + b\sqrt{n''})} \right]^{k-1} \\ &\cong \exp \left\{ - \frac{(k-1) \exp [-(w_A + b\sqrt{n''})^2/2]}{\sqrt{2\pi}(w_A + b\sqrt{n''})} \right\} \end{aligned} \quad (\text{A61})$$

where only the leading term is considered. Hence from (A59), (A60)

and (A61)

$$\begin{aligned}
 -\ln \ln \left(\frac{1-\varepsilon}{P^*} \right)^{\sqrt{2\pi}} + \ln(k-1) \\
 \cong \frac{1}{2}(w_A + b\sqrt{n''})^2 + \ln(w_A + b\sqrt{n''})
 \end{aligned} \tag{A62}$$

Since w_A is bounded and $\ln \sqrt{n''} = o(n'')$ it follows that for large k

$$n'' \cong (2/b^2) \ln(k-1) \cong C \ln k \tag{A63}$$

where C is a proportionality factor. Starting with (A54) and (A56) the same argument gives the same result as (A63) for n' . Hence, by (A58), the same result must hold for n .

ACKNOWLEDGMENT

The authors wish to thank R. B. Murphy, J. W. Tukey, E. L. Kaplan, S. S. Gupta, E. Bleicher, S. Monro, all of Bell Telephone Laboratories, and Prof. R. E. Bechhofer of Cornell University for helpful suggestions and constructive criticism in connection with this paper.

REFERENCES

1. Bechhofer, R. E., and Sobel, M., On a Class of Sequential Multiple Decision Procedures for Ranking Parameters of Koopman-Darmois Populations with Special Reference to Means of Normal Populations, in preparation.
2. Bechhofer, R. E., A Single-Sample Multiple Decision Procedure for Ranking Means of Normal Populations with Known Variances, *Ann. Math. Stat.*, **25**, pp. 16-39, 1954.
3. H. E. Salzer, R. Zucker, and R. Capuano, Table of the Zeros and Weight Factors of the First Twenty Hermite Polynomials, *Journal of Research of the N. B. S.*, **48**, pp. 111-116, 1952.
4. National Bureau of Standards, Tables of the Binomial Probability Distribution, *App. Math. Series* 6, 1950.
5. Ordnance Corps, Tables of the Cumulative Binomial Probabilities, ORDP 20-1, 1952.
6. Harvard Computation Laboratory, Tables of the Cumulative Binomial Probability Distribution, Harvard University Press, Cambridge, 1955.

Bell System Technical Papers Not Published in This Journal

ALLISON, H. W., see Moore, G. E.

ANDERSON, P. W., and TALMAN, J. D.¹

Pressure Broadening of Spectral Lines at General Pressures, Conf. Proc. Breadth of Spectral Lines, pp. 29-61, Oct., 1956.

ARNOLD, S. M.¹

The Growth and Properties of Metal Whiskers, Tech. Proc. Am. Electroplaters Soc., pp. 26-31, 1956.

BALA, V. B., see Geller, S.

BASHKOW, T. R.¹

Effect of Nonlinear Collector Capacitance on Collector Current Rise Time, Trans. I.R.E. PGED, **ED-3**, pp. 167-172, Oct., 1956.

BEACH, A. L., see Thurmond, C. D.

BIRDSALL, H. A.,¹ and GILKENSEN, P. B.³

The Application of Electrical Instruments for Measuring Moisture Contents of Textiles, Am. Dyestuff Reporter, Proc. Am. Assoc. Tex. Chem. and Colorists, **45**, pp. 935-945, Dec. 17, 1956. Committee report of which Messrs. Birdsall and Gilkensen were members.

BRIDGERS, H. E.,¹ and KOLB, E. D.¹

The Distribution Coefficient of Boron in Germanium, J. Chem. Phys., **25**, pp. 648-650, Oct., 1956.

¹ Bell Telephone Laboratories, Inc.

³ Western Electric Company.

BUCK, T. M.,¹ and McKIM, F. S.¹

Depth of Surface Damage Due to Abrasion on Germanium, J. Electrochem. Soc., **103**, pp. 593-597, Nov., 1956.

COMPTON, K. G.¹

Potential Criteria for the Cathodic Protection of Lead Cable Sheath, Corrosion, **12**, pp. 37-44, Nov., 1956.

DAVID, E. E., JR.,¹ and McDONALD, H. S.¹

Note on Pitch-Synchronous Processing of Speech, J. Acous. Soc. Am., **28**, pp. 1261-1266, Nov., 1956.

DICKINSON, D. J.,⁴ POLLAK, H. O.,¹ and WANNIER, G. H.¹

On a Class of Polynomials Orthogonal Over a Denumerable Set, Pacific J. Math., **6**, pp. 239-247, 1956.

FELKER, J. H.¹

Complexity With Reliability, I.R.E. Student Quarterly, **3**, pp. 7-11, Dec., 1956.

GELLER, S.¹

A Set of Effective Coordination Number (12) Radii for the β -Wolfram Structure Elements, Acta Crys., **9**, pp. 885-899, Nov. 10, 1956.

GELLER, S.,¹ and BALA, V. B.¹

Crystallographic Studies of Perovskite-like Compounds. II — Rare Earth Aluminates, Acta Crys., **9**, pp. 1019-1025, Dec. 10, 1956.

GULDNER, W. G.¹

The Application of Vacuum Techniques to Analytical Chemistry, Vakuum-Technik, pp. 159-166, Oct., 1956.

GULDNER, W. G.¹

Tentative Method for Analysis of Carbon in Nickel, A.S.T.M., Chem. Anal. Electronic Nickel (E107-56T), pp. 20-25, Sept. 1956.

¹ Bell Telephone Laboratories, Inc.

⁴ Pennsylvania State University, University Park.

GULDNER, W. G.¹

Tentative Method of Test for Oxygen, Hydrogen and Nitrogen in Nickel, A.S.T.M., Chem. Anal. Electronic Nickel (E107-56T), pp. 26-33, Sept., 1956.

GULDNER, W. G., see Thurmond, C. D.

HAGSTRUM, H. D.¹

Auger Ejection of Electrons from Molybdenum by Noble Gas Ions, Phys. Rev., **104**, pp. 672-683, Nov. 1, 1956.

Auger Ejection of Electrons from Tungsten by Noble Gas Ions, Phys. Rev., **104**, pp. 317-318, Oct. 15, 1956.

Metastable Ions of the Noble Gases, Phys. Rev., **104**, pp. 309-316, Oct. 15, 1956.

HARING, H. E., see Taylor, R. L.

KLEMM, G. H.¹

Automatic Projection Switching for TD-2 Radio System, Commun. and Electronics, **27**, pp. 520-527, Nov., 1956.

KOLB, E. D., see Bridgers, H. E.

KOMPFNER, R.¹

Some Recollections of the Early History of the Traveling Wave Tube, 1956 Yearbook Phys. Soc. London, pp. 30-33, 1956.

KRUSEMEYER, H. J.,¹ and PURSLEY, M. V.¹

Donor Concentration Changes in Oxide Coated Cathodes Due to Changes in Electric Field, J. Appl. Phys., **27**, pp. 1537-1545, Dec., 1956.

LEWIS, H. W.¹

Surface Energies in Superconductors, Phys. Rev., **104**, pp. 942-947, Nov. 15, 1956.

LOVELL, L. CLARICE, see Vogel, F. L., Jr.

¹ Bell Telephone Laboratories, Inc.

MALLINA, R. F.¹

Solderless Wrapped Connections, Trans. I.R.E., **PGT-1**, pp. 12-22, Sept., 1956.

MATTHIAS, B. T.,¹ MILLER, C. E.,¹ and REMEIKA, J. P.¹

Ferroelectricity of Glycine Sulfate, Phys. Rev., **104**, pp. 849-850, Nov. 1, 1956.

MASON, W. P.¹

Internal Friction and Fatigue in Metals at Large Strain Amplitudes, J. Acous. Soc. Am., **28**, pp. 1207-1218, Nov., 1956.

Physical Acoustics and the Properties of Solids, J. Acous. Soc. Am., **28**, pp. 1197-1206, Nov., 1956.

MATREYEK, W., see Winslow, F. H.

MCDONALD, H. S., see David, E. E., Jr.

McKIM, F. S., see Buck, T. M.

McSKIMIN, H. J.¹

Wave Propagation and the Measurement of the Elastic Properties of Liquids and Solids, J. Acous. Soc. Am., **28**, pp. 1228-1232, Nov., 1956.

MILLER, C. E., see Matthias, B. T.

MILLER, R. C.,¹ and SAVAGE, A.¹

Diffusion of Aluminum in Single Crystal Silicon, J. Appl. Phys., **27**, pp. 1430-1432, Dec., 1956.

MONFORTE, F. R., see Van Uitert, L. G.

MOORE, G. E.,¹ and ALLISON, H. W.¹

Emission of Oxide Cathodes Supported on a Ceramic, J. Appl. Phys., **27**, pp. 1316-1321, Nov., 1956.

OSWALD, A. A.¹

Early History of Single Sideband Transmission, Proc. I.R.E., **44**, pp. 1676-1679, Dec., 1956.

¹ Bell Telephone Laboratories, Inc.

PIERCE, J. R.,¹ and WALKER, L. R.¹

Growing Electric Space-Charge Waves, *Phys. Rev.*, **104**, pp. 306-307, Oct. 15, 1956.

POLLAK, H. O., see Dickinson, D. J.

PURSLEY, M. V., see Krusemeyer, H. J.

REMEIKA, J. P., see Matthias, B. T.

RICE, J. W.³

Manufacture of Wire Spring Relays for Communication Switching Systems, *Commun. and Electronics*, **27**, pp. 513-518, Nov., 1956.

ROSE, D. J.¹

The Townsend Ionization Coefficient for Hydrogen and Deuterium, *Phys. Rev.*, **104**, pp. 273-277, Oct. 15, 1956.

SAVAGE, A., see Miller, R. C.

STRUTHERS, J. D.¹

Solubility and Diffusivity of Gold, Iron, and Copper in Silicon, *J. Appl. Phys.*, Letter to the Editor, **27**, p. 1560, Dec., 1956.

SUHL, H., see Walker, L. R.

SWANEKAMP, F. W., see Van Uitert, L. G.

TAYLOR, R. L.,¹ and HARING, H. E.¹

Metal-Semiconductor Capacitor, *J. Electrochem. Soc.*, **103**, pp. 611-613, Nov., 1956.

TALMAN, J. D., see Anderson, P. W.

THURMOND, C. D.,¹ GULDNER, W. G.,¹ and BEACH, A. L.¹

THURMOND, C. D.,¹ GULDNER, W. G.,¹ and BEACH, A. L.¹

Hydrogen and Oxygen in Single-Crystal Germanium as Determined by Vacuum Fusion Gas Analysis, *J. Electrochem. Soc.*, **103**, pp. 603-605, Nov., 1956.

¹ Bell Telephone Laboratories, Inc.

TORREY, MARY N.¹

Quality Control in Electronics, Proc. I.R.E., **44**, pp. 1521-1530, Nov., 1956.

TRUMBORE, F. A.¹

Solid Solubilities and Electrical Properties of Tin in Germanium Single Crystals, J. Electrochem. Soc., **103**, pp. 597-600, Nov., 1956.

VAN UITERT, L. G.,¹ SWANEKAMP, F. W.,¹ and MONFORTE, F. R.¹

Method for Forming Large Ferrite Parts for Microwave Applications, J. Appl. Phys., Letter to the Editor, **27**, pp. 1385-1386, Nov., 1956.

VOGEL, F. L., JR.,¹ and LOVELL, L. CLARICE¹

Dislocation Etch Pits in Silicon Crystals, J. Appl. Phys., **27**, pp. 1413-1415, Dec., 1956.

WALKER, L. R., see Pierce, J. R.

WALKER, L. R.,¹ and SUHL, H.¹

Propagation in Circular Waveguides Filled With Gyromagnetic Material, Trans. I.R.E., **AP-4**, pp. 492-494, July, 1956.

WANNIER, G. H., see Dickinson, D. J.

WEINREICH, G.¹

Acoustodynamic Effects in Semiconductors, Phys. Rev., **104**, pp. 321-324, Oct. 15, 1956.

WERTHEIM, G. K.¹

Carrier Lifetime in Indium Antimonide, Phys. Rev., **104**, pp. 662-664, Nov. 1, 1956.

WINSLOW, F. H.,¹ and MATREYEK, W.¹

Pyrolysis of Cross-Linked Styrene Polymers, J. Poly. Sci., **22**, pp. 315-324, Nov., 1956.

¹ Bell Telephone Laboratories, Inc.

Recent Monographs of Bell System Technical Papers Not Published in This Journal*

ANDERSON, O. L.

Effect of Pressure on Glass Structure, Monograph 2666.

BEMSKI, G.

Quenched-In Recombination Centers in Silicon, Monograph 2681.

BRIDGERS, H. E.

p-n Junctions in Semiconductors by Variation of Crystal Growth Parameters, Monograph 2668.

BROWN, W. L., see Montgomery, H. C.

CAMPBELL, MARY E., see Luke, C. L.

CARLITZ, L., and RIORDAN, J.

The Number of Labeled Two-Terminal Series-Parallel Networks, Monograph 2667.

CHASE, F. H.

Power Regulation by Semiconductors, Monograph 2685.

DAVID, E. E., JR.

Naturalness and Distortion in Speech-Processing Devices, Monograph 2687.

DODGE, H. F., and TORREY, Miss M. N.

A Check Inspection and Demerit Rating Plan, Monograph 2669.

* Copies of these monographs may be obtained on request to the Publication Department, Bell Telephone Laboratories, Inc., 463 West Street, New York 14, Y. Y. The numbers of the monographs should be given in all requests.

EDER, M. J., see Veloric, H. S.

FISHER, J. R., and POTTER, J. F.

Apparent Density for Evaluating the Physical Structure of Steatite, Monograph 2688.

FULLER, C. S., see Reiss, H.

GELLER, S., and GILLES, M. A.

Gadolinium Orthoferrite: Crystal Structure and Magnetic Properties, Monograph 2670.

GILLES, M. A., see Geller, S.

GOHN, G. R.

Hardness Conversion Table for Copper-Beryllium Alloy Strip, Monograph 2665.

HARROWER, G. A.

Dependence of Electron Reflection on Contamination of Reflecting Surface, Monograph 2689.

HARROWER, G. A.

Energy Spectra of Secondary Electrons from Mo and W for Low Primary Energies, Monograph 2690.

HØVGAARD, O. M.

Capability of Sealed Contact Relays, Monograph 2697.

KNAPP, H. M.

Design Features of Bell System Wire Spring Relays, Monograph 2693.

LEWIS, H. W.

Two-Fluid Model of an "Energy-Gap" Superconductor, Monograph 2671.

LUKE, C. L., and CAMPBELL, MARY E.

Photometric Determination of Germanium and Tin With Phenyl-fluorone, Monograph 2608.

MANLEY, J. M., and ROWE, H. E.

Some General Properties of Nonlinear Elements. I — General Energy Relations, Monograph 2672.

MCLEAN, D. A., and POWER, F. S.

Tantalum Solid Electrolytic Capacitors, Monograph 2673.

MCMAHON, W.

Dielectrics by Solidifying Certain Organic Compounds in Electric or Magnetic Fields, Monograph 2694.

MONTGOMERY, H. C., and BROWN, W. L.

Field-Induced Conductivity Changes in Germanium, Monograph 2695.

MOORE, E. F., and SHANNON, C. E.

Reliable Circuits Using Less Reliable Relays, Monograph 2696.

PIETRUSZKIEWICZ, A. J., see Reiss, H.

POTTER, J. F., see Fisher, J. R.

POWER, F. S., see McLean, D. A.

PRINCE, M. B., see Veloric, H. S.

REISS, H.

Refined Theory of Ion Pairing, Monograph 2698.

REISS, H., FULLER, C. S., and PIETRUSZKIEWICZ, A. J.

Solubility of Lithium in Doped and Undoped Silicon, Monograph 2702.

RIESS, H.

Theory of Ionization of Hydrogen and Lithium in Silicon and Germanium, Monograph 2700.

REMEIKA, J. P.

Growth of Single Crystal Rare-Earth Orthoferrites and Related Compounds, Monograph 2699.

RICHARDS, A. P., see Snoke, L. R.

RIORDAN, J., see Carlitz, L.

ROWE, H. E., see Manley, J. M.

SHANNON, C. E., see Moore, E. F.

SHULMAN, R. G., and WYLUDA, B. J.

Trapping Center Properties of Germanium, Monograph 2674.

SNOKE, L. R., and RICHARDS, A. P.

Marine Borer Attack on Lead Cable Sheath, Monograph 2675.

TORREY, Miss M. N., see Dodge, H. F.

UHLIR, A., JR.

Two-Terminal p-n Junction Devices for Frequency Conversion and Computation, Monograph 2704.

VAN UITERT, L. G.

Nickel Copper Ferrites for Microwave Applications, Monograph 2676.

VELORIC, H. S., PRINCE, M. B., and EDER, M. J.

Avalanche Breakdown Voltage in Silicon Diffused p-n Junctions, Monograph 2705.

WEINREICH, G.

Transit Time Transistor, Monograph 2706.

WERNICK, J. H.

Diffusivities in Liquid Metals by Temperature-Gradient Zone Melting, Monograph 2677.

WOLFF, P. A.

Theory of Plasma Resonance, Monograph 2707.

WYLUDA, B. J., see Shulman, R. G.

Contributors to This Issue

RICHARD C. BOYD, B.S., Northwestern University, 1946; B.S.E.E., University of Michigan, 1947; M.S.E.E., University of Michigan, 1948; Bell Telephone Laboratories, 1948-. After completion of the Laboratories Communication Development Training program in 1950, Mr. Boyd was concerned with transmission engineering and systems studies. He was a supervisor during the exploratory trial of experimental P1 carrier in 1953 and 1954. He is now responsible for transmission engineering of P1 carrier for rural subscriber use, and for adaptation of P and N1 carrier to exchange trunk use. He is a member of Tau Beta Pi and Phi Kappa Phi.

ROBERT W. DAWSON, Newark College of Engineering; Rutgers University; Bell Telephone Laboratories, 1941-. All of Mr. Dawson's work has been with the Radio Research Department. Together with A. C. Beck, he was concerned with conductivity measurements at microwave frequencies, and they were co-authors of an article on this subject. Mr. Dawson, a member of the Institute of Radio Engineers, worked on the Manhattan Project at Los Alamos while serving in the Army from 1942 to 1946.

GEORGE FEHER, B.S. in Engineering Physics, University of California, 1950; M.S.E.E., University of California, 1952; Ph.D. in Physics, University of California, 1954; Bell Telephone Laboratories, 1954-. Dr. Feher has been a member of the Semiconductor Research Department since joining the Laboratories. He has been engaged particularly in studies of electron spin resonance absorption, and is the author of several articles on this subject. He is a member of the American Physical Society, the Institute of Radio Engineers, and Sigma Xi.

JOHN D. HOWARD, B.E.E., University of Louisville, 1947; Southern Bell Telephone Company, 1947-. After a year and a half in the Plant Department at Southern Bell, Mr. Howard joined the Engineering Department and remained there until late 1952. At that time he was called to the O. & E. Department of the A. T. & T. Co. where he was concerned with exchange transmission. After completing this assignment he

returned to Southern Bell where he is now Exchange Transmission Engineer. Mr. Howard is a member of the A.I.E.E.

MARILYN J. HUYETT, A.B., Susquehanna University, 1954; Bell Telephone Laboratories, 1954-. Since joining the Laboratories, Miss Huyett has been concerned with statistical research for the reliability group at Allentown, where she has had a great deal of experience with IBM computing machines. She is a member of the American Statistical Association. While in college, she received the Stine Mathematical Prize, an award for proficiency in mathematics.

JOHN E. KARLIN, B.A., University of Cape Town, 1938; M.A., University of Cape Town, 1939; Ph.D., University of Chicago, 1942; Bell Telephone Laboratories, 1945-. Dr. Karlin was engaged in psycho-acoustic research from 1945 to 1947. From 1947 until the present he has been concerned with user preference research, and since 1952 he has been in charge of the group doing this research. During the war years, 1942-1945, Dr. Karlin was at Harvard University, engaged in communications research on military projects. He is a member of the I.R.E., the Acoustical Society of America, the American Psychological Association and Sigma Xi.

STUART P. LLOYD, S.B., 1943, University of Chicago; M.S., 1949 and Ph.D., 1951, University of Illinois. Bell Telephone Laboratories, 1952-. Engaged in work relating to probability theory and information theory. Member, Institute for Advanced Study, Princeton, 1951-52. Member of American Mathematical Society, Institute of Mathematical Statistics, American Physical Society, A.A.A.S., Phi Kappa Phi, Sigma Xi and Phi Beta Kappa.

D. W. MCCALL, B.S., University of Wichita, 1950; M.S., 1951 and Ph.D., 1953, University of Illinois; Bell Telephone Laboratories, 1953-. Dr. McCall is engaged in fundamental studies of the properties of dielectrics. He is a member of the American Chemical Society, the American Physical Society, Sigma Xi, and Phi Lambda Upsilon.

LUDWIG PEDERSEN, Christiania Technical School (Norway), 1919; Western Electric International Co. (Oslo), 1919-20; Western Electric Co., 1920-25; Western Electric Co., Field Engineering Force, 1944-45; Bell Telephone Laboratories, 1925-. He has been concerned with circuit design for machine switching systems, development of telegraph equip-

ment, design of transmission equipment for the armed forces, and toll transmission systems. He served as a technical observer with the U. S. Army in the European theater. He is now Systems Development Engineer at the Merrimack Valley Laboratory with responsibility for voice frequency, broadband carrier, short haul carrier and rural carrier systems. Member of the A.I.E.E.

JOHN R. PIERCE, B.S., 1933, M.S., 1934 and Ph.D., 1936, California Institute of Technology; Bell Telephone Laboratories, 1936-. Director of Research in Electrical Communications at Bell Telephone Laboratories. He has specialized in the development of electron tubes, microwave research, electronic devices for military applications, and communications circuits. Dr. Pierce has been granted 55 patents and is the author of three books. For his research leading to the development of the beam traveling wave tube, he was awarded the 1947 Morris Liebmman Memorial Prize of the I. R. E. He was voted the "Outstanding Young Electrical Engineer of 1942" by Eta Kappa Nu. Member of National Academy of Sciences, British Interplanetary Society, A.I.E.E., Sigma Xi, Tau Beta Pi and Eta Kappa Nu. Fellow of American Physical Society and I.R.E.

JOHN H. ROWEN, B.E.E., 1948 and M.Sc., 1951, Ohio State University. Mr. Rowen worked at the Antenna Laboratory of the Ohio State University Research Foundation in Columbus, O. from 1948 to 1951. He joined Bell Telephone Laboratories in 1951, shortly after being awarded his master's degree. Since then Mr. Rowen has specialized in the applied physics of solids, and the development of microwave ferrite devices. He is a member of the Solid State Device Development Department at the Murray Hill Laboratory. Mr. Rowen is a member of the Institute of Radio Engineers and of Eta Kappa Nu.

HAROLD SEIDEL, B.E.E., 1943, College of the City of New York; M.E.E., 1947 and D.E.E., 1954, Polytechnic Institute of Brooklyn; Microwave Research Institute of Polytechnic Institute of Brooklyn, 1947; Arma Corp., 1947-48; Federal Telecommunications Laboratories, 1948-53; Bell Telephone Laboratories, 1953-. He has been concerned with general electromagnetic problems, especially regarding wave-guide applications, and with analysis of microwave ferrite devices. Member of Sigma Xi and I.R.E.

MILTON SOBEL, B. S., 1940, College of the City of New York; M.A., 1946 and Ph.D., 1951, Columbia University; U. S. Census Bureau, stat-

istician, 1940-41; U. S. Army War College, statistician, 1942-44; Columbia University, department of mathematics, 1946-50; Wayne University, assistant professor of mathematics, 1950-52; Columbia University, visiting lecturer, 1952; Cornell University, fundamental research in mathematical statistics, 1952-54; Bell Telephone Laboratories, 1954-. He has been engaged in fundamental research on life testing and reliability problems with special application to transistors. Consultant on many Bell Laboratories projects. Member of Institute of Mathematical Statistics, American Statistical Association and Sigma Xi.

WILHELM VON AULOCK, Dipl. Ing., Technische Hochschule, Berlin, 1937; Dr. Ing., Technische Hochschule Stuttgart, 1953; Dr. von Aulock worked in Berlin from 1938 to 1942 for the A.E.G., Kabelwerk; in Gotenhafen (Gdynia) for the Torpedoversuchsanstalt from 1942 to 1945; and for the United States Navy Bureau of Ships in Washington from 1947 to 1953, where he was involved in work on torpedo counter-measures and studies of electromagnetic induction fields in sea water. He joined Bell Telephone Laboratories in 1954. Since then he has been principally engaged in analytical and experimental studies of phase shift and loss characteristics of ferrite-loaded waveguides, and the application of these properties. He is an associate member of the Institute of Radio Engineers.

THE BELL SYSTEM

Technical Journal

VOTED TO THE SCIENTIFIC AND ENGINEERING
PECTS OF ELECTRICAL COMMUNICATION

VOLUME XXXVI

MAY 1957

NUMBER 3

Radio Propagation Fundamentals KENNETH BULLINGTON 593

A Reflection Theory for Propagation Beyond the Horizon

H. T. FRIIS, A. B. CRAWFORD, D. C. HOGG 627

Interchannel Interference Due to Klystron Pulling

H. E. CURTIS AND S. O. RICE 645

Instantaneous Companding of Quantized Signals

BERNARD SMITH 653

An Electrically Operated Hydraulic Control Valve

J. W. SCHAEFER 711

Strength Requirements for Round Conduit

G. F. WEISSMANN 737

Cold Cathode Gas Tubes for Telephone Switching Systems

M. A. TOWNSEND 755

Activation of Electrical Contacts by Organic Vapors

L. H. GERMER AND J. L. SMITH 769

Bell System Technical Papers Not Published in This Journal 813

Recent Bell System Monographs

KANSAS CITY, MO.
PUBLIC LIBRARY

823

Contributors to This Issue

JUN 7 1957

828

THE BELL SYSTEM TECHNICAL JOURNAL

ADVISORY BOARD

A. B. GOETZE, *President, Western Electric Company*

M. J. KELLY, *President, Bell Telephone Laboratories*

E. J. MCNEELY, *Executive Vice President, American Telephone and Telegraph Company*

EDITORIAL COMMITTEE

B. MCMILLAN, *Chairman*

S. E. BRILLHART

E. I. GREEN

A. J. BUSCH

R. K. HONAMAN

L. R. COOK

H. R. HUNTLEY

A. C. DICKIESON

F. R. LACK

R. L. DIETZOLD

J. R. PIERCE

K. E. GOULD

G. N. THAYER

EDITORIAL STAFF

W. D. BULLOCH, *Editor*

R. L. SHEPHERD, *Production Editor*

T. N. POPE, *Circulation Manager*

THE BELL SYSTEM TECHNICAL JOURNAL is published six times a year by the American Telephone and Telegraph Company, 195 Broadway, New York 7, N. Y. F. R. Kappel, President; S. Whitney Landon, Secretary; John J. Scanlon, Treasurer. Subscriptions are accepted at \$5.00 per year. Single copies \$1.25 each. Foreign postage is 65 cents per year or 11 cents per copy. Printed in U. S. A.

THE BELL SYSTEM TECHNICAL JOURNAL

VOLUME XXXVI

MAY 1957

NUMBER 3

Copyright 1957, American Telephone and Telegraph Company

Radio Propagation Fundamentals*

By KENNETH BULLINGTON

(Manuscript received June 21, 1956)

The engineering of radio systems requires an estimate of the power loss between the transmitter and the receiver. Such estimates are affected by many factors, including reflections, fading, refraction in the atmosphere, and diffraction over the earth's surface.

In this paper, radio transmission theory and experiment in all frequency bands of current interest are summarized. Ground wave and sky wave transmission are included, and both line of sight and beyond horizon transmission are considered. The principal emphasis is placed on quantitative charts that are useful for engineering purposes.

I. INTRODUCTION

The power radiated from a transmitting antenna is ordinarily spread over a relatively large area. As a result the power available at most receiving antennas is only a small fraction of the radiated power. This ratio of radiated power to received power is called the radio transmission loss and its magnitude in some cases may be as large as 10^{15} to 10^{20} (150 to 200 decibels).

The transmission loss between the transmitting and receiving antennas determines whether the received signal will be useful. Each radio

* This paper has been prepared for use in a proposed "Antenna Handbook" to be published by McGraw-Hill.

system has a maximum allowable transmission loss which, if exceeded, results in either poor quality or poor reliability. Reasonably accurate predictions of transmission loss can be made on paths that approximate the ideals of either free space or plane earth. On many paths of interest, however, the path geometry or atmospheric conditions differ so much from the basic assumptions that absolute accuracy cannot be expected; nevertheless, worthwhile results can be obtained by using two or more different methods of analysis to "box in" the answer.

The basic concept in estimating radio transmission loss is the loss expected in free space; that is, in a region free of all objects that might absorb or reflect radio energy. This concept is essentially the inverse square law in optics applied to radio transmission. For a one wavelength separation between nondirective (isotropic) antennas, the free space loss is 22 db and it increases by 6 db each time the distance is doubled. The free space transmission ratio at a distance d is given by:

$$\frac{P_r}{P_t} = \left(\frac{\lambda}{4\pi d} \right)^2 g_t g_r \quad (1a)$$

where:

$$\left. \begin{array}{l} P_r = \text{received power} \\ P_t = \text{radiated power} \end{array} \right\} \text{measured in same units}$$

$$\lambda = \text{wavelength in same units as } d$$

$$g_t \text{ (or } g_r) = \text{power gain of transmitting (or receiving) antenna}$$

The power gain of an ideal isotropic antenna that radiates power uniformly in all directions is unity by definition. A small doublet whose over-all physical length is short compared with one-half wavelength has a gain of $g = 1.5$ (1.76 decibels) and a one-half wave dipole has a gain of 2.15 decibels in the direction of maximum radiation. A nomogram for the free space transmission loss between isotropic antennas is given in Fig. 1.

When antenna dimensions are large compared with the wavelength, a more convenient form of the free space ratio is¹

$$\frac{P_r}{P_t} = \frac{A_t A_r}{(\lambda d)^2} \quad (1b)$$

where $A_{t,r}$ = effective area of transmitting or receiving antennas.

Another form of expressing free space transmission is the concept of

the free space field intensity E_0 which is given by:

$$E_0 = \frac{\sqrt{30P_t g_t}}{d} \text{ volts per meter} \quad (2)$$

where d is in meters and P_t in watts.

The use of the field intensity concept is frequently more convenient than the transmission loss concept at frequencies below about 30 mc,

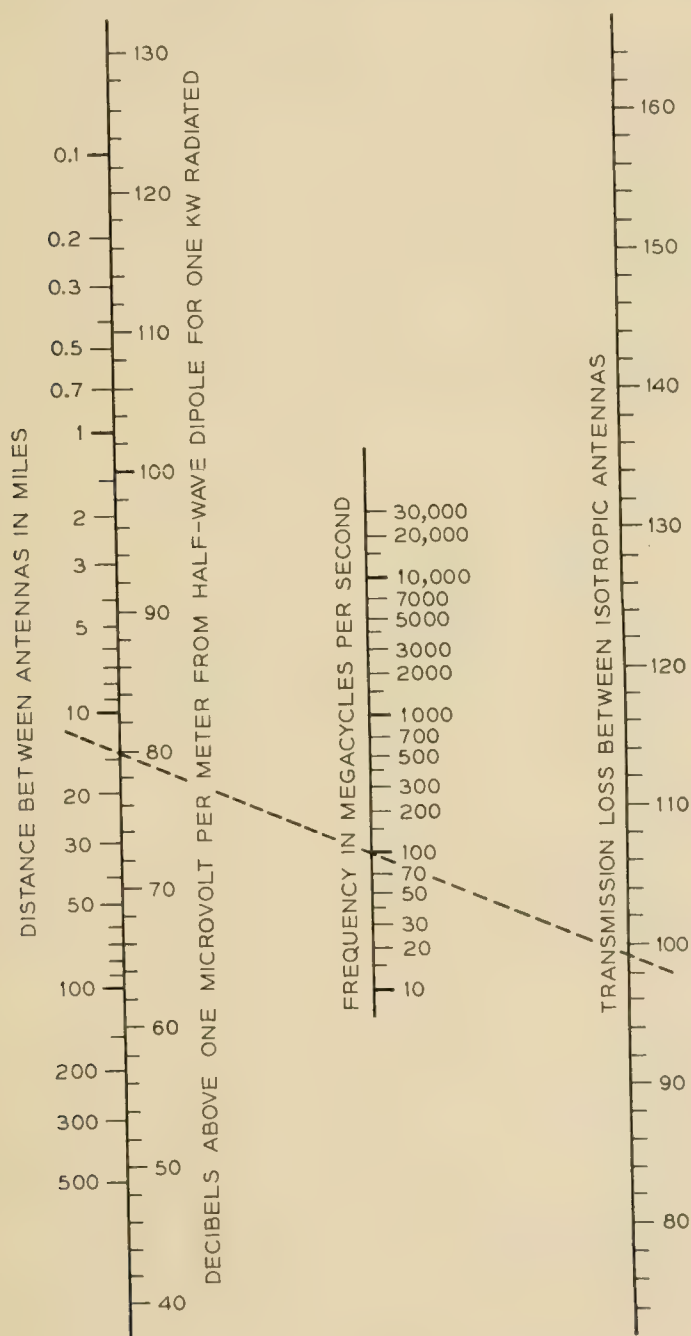


FIG. 1 — Free space transmission.

where external noise is generally controlling and where antenna dimensions and heights are comparable to or less than a wavelength. The free space field intensity is independent of frequency and its magnitude for one kilowatt radiated from a half-wave dipole is shown on the left hand scale on Fig. 1.

The concept of free space transmission assumes that the atmosphere is perfectly uniform and nonabsorbing and that the earth is either infinitely far away or its reflection coefficient is negligible. In practice, the modifying effects of the earth, the atmosphere and the ionosphere need to be considered. Both theoretical and experimental values for these effects are described in the following sections.

II. TRANSMISSION WITHIN LINE OF SIGHT

The presence of the ground modifies the generation and the propagation of radio waves so that the received power or field intensity is ordinarily less than would be expected in free space.² The effect of plane earth on the propagation of radio waves is given by

$$\frac{E}{E_0} = \underbrace{1}_{\text{Direct Wave}} + \underbrace{Re^{j\Delta}}_{\text{Reflected Wave}} + \underbrace{(1 - R)Ae^{j\Delta}}_{\text{"Surface Wave" Induction Field and Secondary Effects of the Ground}} + \dots \quad (3)$$

where

R = reflection coefficient of the ground

A = "surface wave" attenuation factor

$$\Delta = \frac{4\pi h_1 h_2}{\lambda d}$$

$h_{1,2}$ = antenna heights measured in same units as the wavelength and distance

The parameters R and A vary with both polarization and the electrical constants of the ground. In addition, the term "surface wave" has led to considerable confusion since it has been used in the literature to stand for entirely different concepts. These factors are discussed more completely in Section IV. However, the important point to note in this section is that considerable simplification is possible in most practical cases, and that the variations with polarization and ground constants

and the confusion about the surface wave can often be neglected. For near grazing paths, R is approximately equal to -1 and the factor A can be neglected as long as both antennas are elevated more than a wavelength above the ground (or more than 5–10 wavelengths above sea water). Under these conditions the effect of the earth is independent of polarization and ground constants and (3) reduces to

$$\left| \frac{E}{E_0} \right| = \sqrt{\frac{P_r}{P_0}} = 2 \sin \frac{\Delta}{2} = 2 \sin \frac{2\pi h_1 h_2}{\lambda d} \quad (4)$$

where P_0 is the received power expected in free space.

The above expression is the sum of the direct and ground reflected rays and shows the lobe structure of the signal as it oscillates around the free space value. In most radio applications (except air to ground) the principal interest is in the lower part of the first lobe; that is, where $\Delta/2 < \pi/4$. In this case, $\sin \Delta/2 \approx \Delta/2$ and the transmission loss over plane earth is given by:

$$\begin{aligned} \frac{P_r}{P_t} &= \left(\frac{\lambda}{4\pi d} \right)^2 \left(\frac{4\pi h_1 h_2}{\lambda d} \right)^2 g_t g_r \\ &= \left(\frac{h_1 h_2}{d^2} \right)^2 g_t g_r \end{aligned} \quad (5)$$

It will be noted that this relation is independent of frequency and it is shown in decibels in Fig. 2 for isotropic antennas. Fig. 2 is not valid when the indicated transmission loss is less than the free space loss shown in Fig. 1, because this means that Δ is too large for this approximation.

Although the transmission loss shown in (5) and in Fig. 2 has been derived from optical concepts that are not strictly valid for antenna heights less than a few wavelengths, approximate results can be obtained for lower heights by using h_1 (or h_2) as the larger of either the actual antenna height or the minimum effective antenna height shown in Fig. 3. The concept of minimum effective antenna height is discussed further in Section IV. The error that can result from the use of this artifice does not exceed ± 3 db and occurs where the actual antenna height is approximately equal to the minimum effective antenna height.

The sine function in (4) shows that the received field intensity oscillates around the free space value as the antenna heights are increased. The first maximum occurs when the difference between the direct and ground reflected waves is a half wavelength. The signal maxima have a magnitude $1 + |R|$ and the signal minima have a magnitude of $1 - |R|$.

Frequently the amount of clearance (or obstruction) is described in terms of Fresnel zones. All points from which a wave could be reflected

with a path difference of one-half wavelength form the boundary of the first Fresnel zone; similarly, the boundary of the n^{th} Fresnel zone consists of all points from which the path difference is $n/2$ wavelengths. The n^{th} Fresnel zone clearance H_n at any distance d_1 is given by:

$$H_n = \sqrt{\frac{n\lambda d_1(d - d_1)}{d}} \quad (6)$$

Although the reflection coefficient is very nearly equal to -1 for

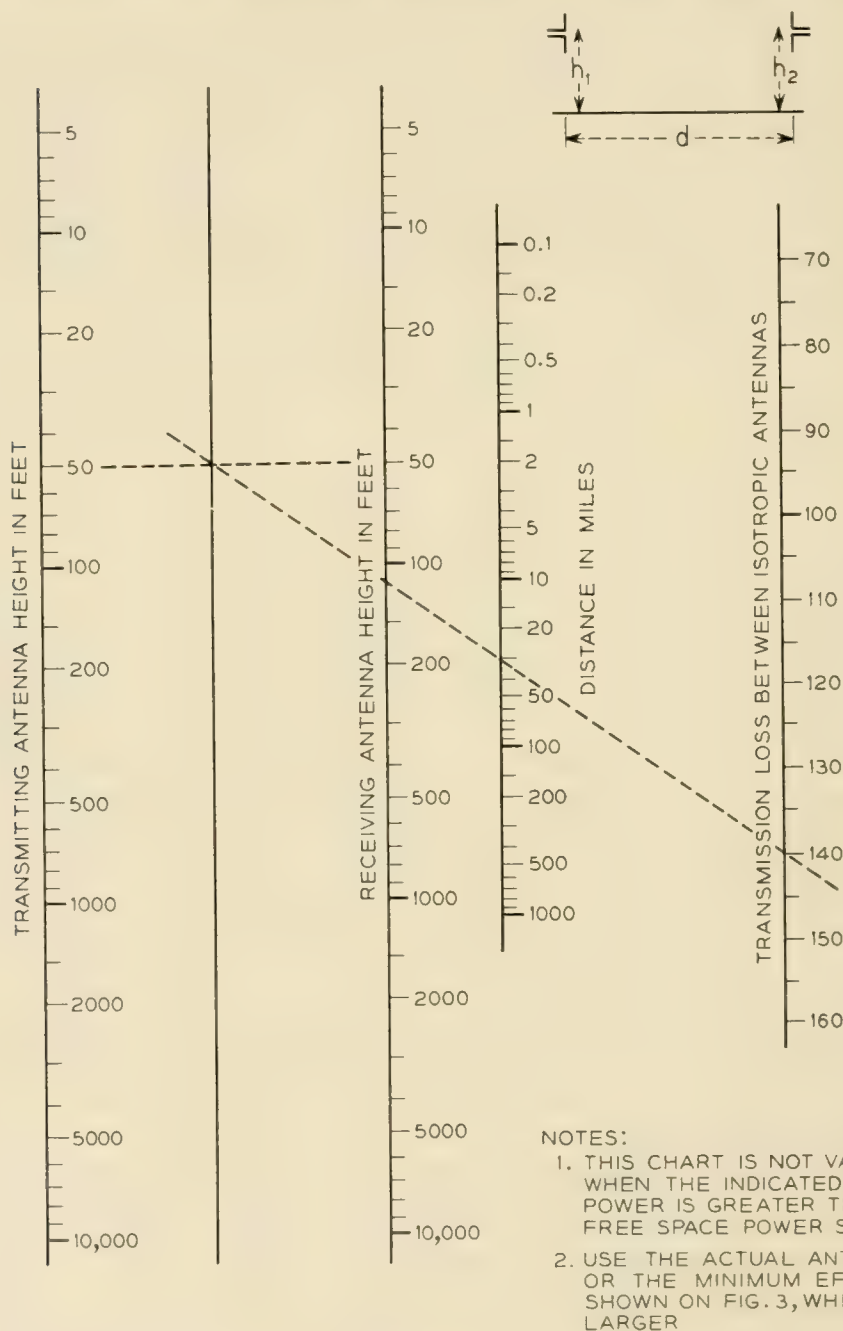


FIG. 2 — Transmission loss over plane earth.

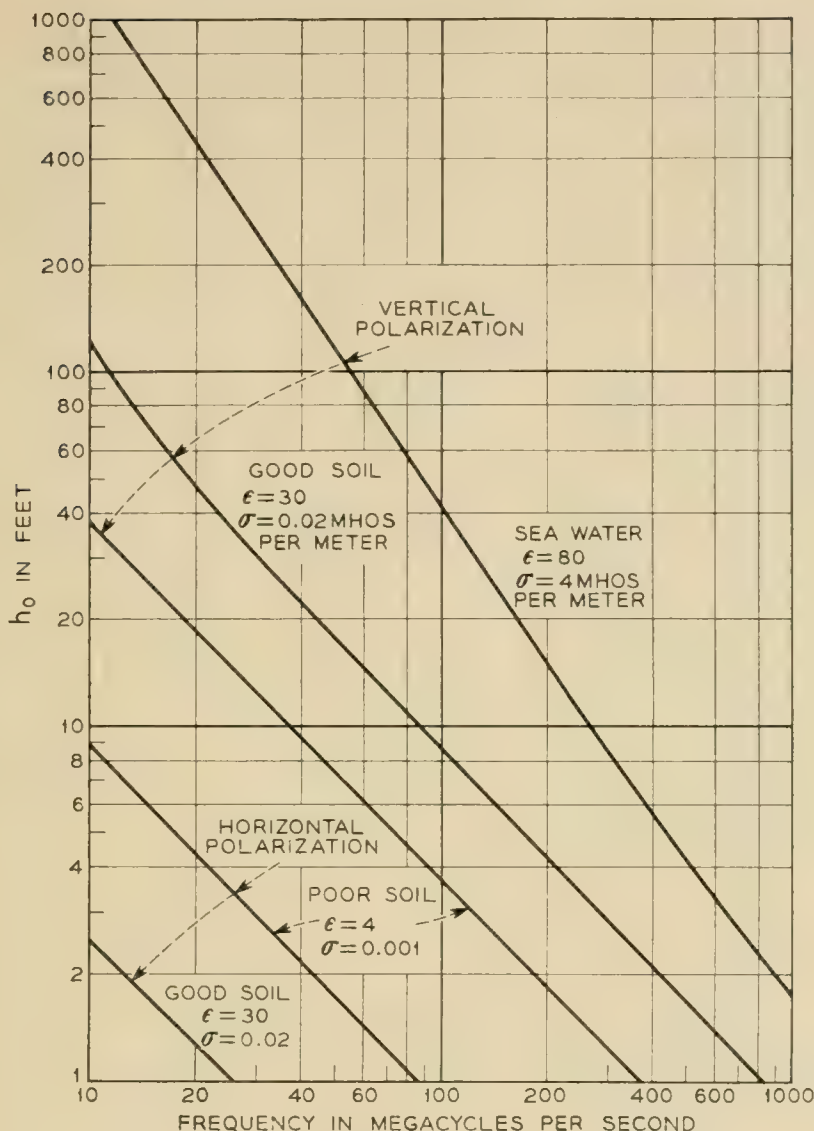


FIG. 3 — Minimum effective antenna height.

grazing angles over smooth surfaces, its magnitude may be less than unity when the terrain is rough. The classical Rayleigh criterion of roughness indicates that specular reflection occurs when the phase deviations are less than about $\pm(\pi/2)$ and that the reflection coefficient will be substantially less than unity when the phase deviations are greater than $\pm(\pi/2)$. In most cases this theoretical boundary between specular and diffuse reflection occurs when the variations in terrain exceed $\frac{1}{8}$ to $\frac{1}{4}$ of the first Fresnel zone clearance. Experimental results with microwave transmission have shown that most practical paths are "rough" and ordinarily have a reflection coefficient in the range of 0.2–0.4. In addition, experience has shown that the reflection coefficient is a statistical problem and cannot be predicted accurately from the path profile.³

Fading Phenomena

Variations in signal level with time are caused by changing atmospheric conditions. The severity of the fading usually increases as either the frequency or path length increases. Fading cannot be predicted accurately but it is important to distinguish between two general types: (1) inverse bending and (2) multipath effects. The latter includes the fading caused by interference between direct and ground reflected waves as well as interference between two or more separate paths in the atmosphere. Ordinarily, fading is a temporary diversion of energy to some other than the desired location; fading caused by absorption of energy is discussed in a later paragraph.

The path of a radio wave is not a straight line except for the ideal case of a uniform atmosphere. The transmission path may be bent up or down depending on atmospheric conditions. This bending may either increase or decrease the effective path clearance and inverse bending may have the effect of transforming a line of sight path into an obstructed one. This type of fading may last for several hours. The frequency of its occurrence and its depth can be reduced by increasing the path clearance, particularly in the middle of the path.

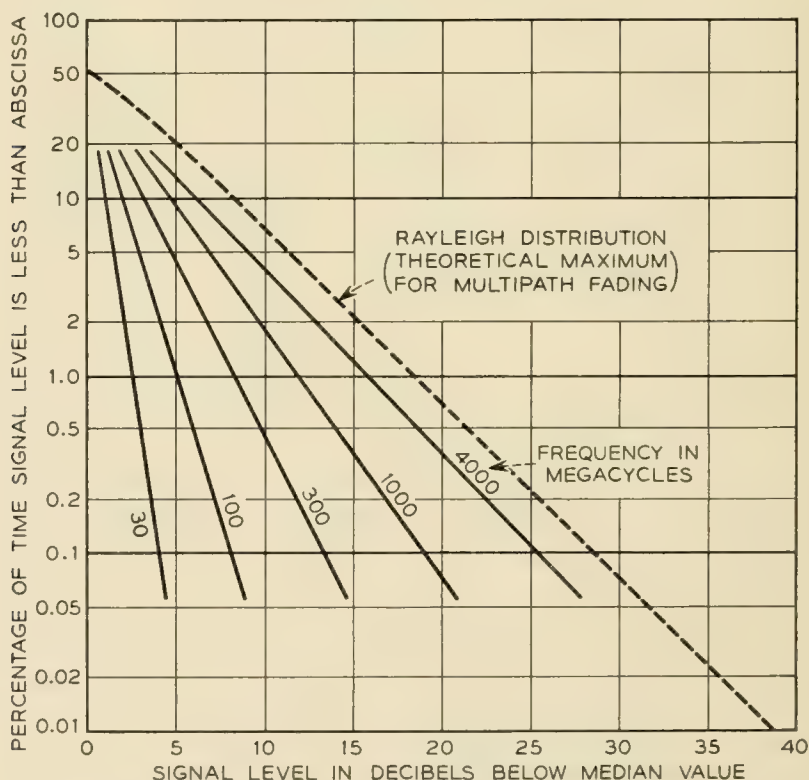


FIG. 4 — Typical fading characteristics in the worst month on 30 to 40 mile line-of-sight paths with 50 to 100 foot clearance.

Severe fading may occur over water or on other smooth paths because the phase difference between the direct and reflected rays varies with atmospheric conditions. The result is that the two rays sometimes add and sometimes tend to cancel. This type of fading can be minimized, if the terrain permits, by locating one end of the circuit high while the other end is very low. In this way the point of reflection is placed near the low antenna and the phase difference between direct and reflected rays is kept relatively steady.

Most of the fading that occurs on "rough" paths with adequate clearance is the result of interference between two or more rays traveling slightly different routes in the atmosphere. This multipath type of fading is relatively independent of path clearance and its extreme condition approaches the Rayleigh distribution. In the Rayleigh distribution, the probability that the instantaneous value of the field is greater than the value R is $\exp [-(R/R_0)]$, where R_0 is the rms value.

Representative values of fading on a path with adequate clearance are shown on Fig. 4. After the multipath fading has reached the Rayleigh distribution, a further increase in either distance or frequency increases the number of fades of a given depth but decreases the duration so that the product is the constant indicated by the Rayleigh distribution.

Miscellaneous Effects

The remainder of this Section describes some miscellaneous effects of line of sight transmission that may be important at frequencies above about 1,000 mc. These effects include variation in angles of arrival, maximum useful antenna gain, useful bandwidth, the use of frequency or space diversity, and atmospheric absorption.

On line of sight paths with adequate clearance some components of the signal may arrive with variations in angle of arrival of as much as $\frac{1}{2}^\circ$ to $\frac{3}{4}^\circ$ in the vertical plane, but the variations in the horizontal plane are less than 0.1° .^{4, 5} Consequently, if antennas with beamwidths less than about 0.5° are used, there may occasionally be some loss in received signal because most of the incoming energy arrives outside the antenna beamwidth. Signal variations due to this effect are usually small compared with the multipath fading.

Multipath fading is selective fading and it limits both the maximum useful bandwidth and the frequency separation needed for adequate frequency diversity. For 40-db antennas on a 30-mile path the fading on frequencies separated by 100–200 mc is essentially uncorrelated regardless of the absolute frequency. With less directive antennas, uncorrelated fading can occur at frequencies separated by less than 100 mc.^{6, 7}

Larger antennas (more narrow beamwidths) will decrease the fast multipath fading and widen the frequency separation between uncorrelated fading but at the risk of increasing the long term fading associated with the variations in the angle of arrival.

Optimum space diversity, when ground reflections are controlling, requires that the separation between antennas be sufficient to place one antenna on a field intensity maximum while the other is in a field intensity minimum. In practice, the best spacing is usually not known because the principal fading is caused by multipath variations in the atmosphere. However, adequate diversity can usually be achieved with a vertical separation of 100–200 wavelengths.

At frequencies above 5,000–10,000 mc, the presence of rain, snow, or fog introduces an absorption in the atmosphere which depends on the amount of moisture and on the frequency. During a rain of cloud burst proportions the attenuation at 10,000 mc may reach 5 db per mile and at 25,000 mc it may be in excess of 25 db per mile.⁸ In addition to the effect of rainfall some selective absorption may result from the oxygen and water vapor in the atmosphere. The first absorption peak due to water vapor occurs at about 24,000 mc and the first absorption peak for oxygen occurs at about 60,000 mc.

III. TROPOSPHERIC TRANSMISSION BEYOND LINE OF SIGHT

A basic characteristic of electromagnetic waves is that the energy is propagated in a direction perpendicular to the surface of uniform phase. Radio waves travel in a straight line only as long as the phase front is plane and is infinite in extent.

Energy can be transmitted beyond the horizon by three principal methods: reflection, refraction and diffraction. Reflection and refraction are associated with either sudden or gradual changes in the direction of the phase front, while diffraction is an edge effect that occurs because the phase surface is not infinite. When the resulting phase front at the receiving antenna is irregular in either amplitude or position, the distinctions between reflection, refraction, and diffraction tend to break down. In this case the energy is said to be scattered. Scattering is frequently pictured as a result of irregular reflections although irregular refraction plus diffraction may be equally important.

The following paragraphs describe first the theories of refraction and of diffraction over a smooth sphere and a knife edge. This is followed by empirical data derived from experimental results on the transmission to points far beyond the horizon, on the effects of hills and trees, and on fading phenomena.

Refraction

The dielectric constant of the atmosphere normally decreases gradually with increasing altitude. The result is that the velocity of transmission increases with the height above the ground and, on the average, the radio energy is bent or refracted toward the earth. As long as the change in dielectric constant is linear with height, the net effect of refraction is the same as if the radio waves continued to travel in a straight line but over an earth whose modified radius is:

$$ka = \frac{a}{1 + \frac{a}{2} \frac{d\epsilon}{dh}} \quad (7)$$

where

a = true radius of earth

$\frac{d\epsilon}{dh}$ = rate of change of dielectric constant with height

Under certain atmospheric conditions the dielectric constant may increase ($0 < k < 1$) over a reasonable height, thereby causing the radio waves in this region to bend away from the earth. This is the cause of the inverse bending type of fading mentioned in the preceding section. It is sometimes called substandard refraction. Since the earth's radius is about 2.1×10^7 feet, a decrease in dielectric constant of only 2.4×10^{-8} per foot of height results in a value of $k = \frac{4}{3}$, which is commonly assumed to be a good average value.⁹ When the dielectric constant decreases about four times as rapidly (or by about 10^{-7} per foot of height), the value of $k = \infty$. Under such a condition, as far as radio propagation is concerned, the earth can then be considered flat, since any ray that starts parallel to the earth will remain parallel.

When the dielectric constant decreases more rapidly than 10^{-7} per foot of height, radio waves that are radiated parallel to, or at an angle above the earth's surface, may be bent downward sufficiently to be reflected from the earth. After reflection the ray is again bent toward the earth, and the path of a typical ray is similar to the path of a bouncing tennis ball. The radio energy appears to be trapped in a duct or waveguide between the earth and the maximum height of the radio path. This phenomenon is variously known as trapping, duct transmission, anomalous propagation, or guided propagation.^{10, 11} It will be noted that in this case the path of a typical guided wave is similar in form to the path of sky waves, which are lower-frequency waves trapped between the

earth and the ionosphere. However, there is little or no similarity between the virtual heights, the critical frequencies, or the causes of refraction in the two cases.

Duct transmission is important because it can cause long distance interference with another station operating on the same frequency; however, it does not occur often enough nor can its occurrence be predicted with enough accuracy to make it useful for radio services requiring high reliability.

Diffraction Over a Smooth Spherical Earth and Ridges

Radio waves are also transmitted around the earth by the phenomenon of diffraction. Diffraction is a fundamental property of wave motion, and in optics it is the correction to apply to geometrical optics (ray theory) to obtain the more accurate wave optics. In other words, all shadows are somewhat "fuzzy" on the edges and the transition from "light" to "dark" areas is gradual, rather than infinitely sharp. Our common experience is that light travels in straight lines and that shadows are sharp, but this is only because the diffraction effects for these very short wavelengths are too small to be noticed without the aid of special laboratory equipment. The order of magnitude of the diffraction at radio frequencies may be obtained by recalling that a 1,000-mc radio wave has about the same wavelength as a 1,000-cycle sound wave in air, so that these two types of waves may be expected to bend around absorbing obstacles with approximately equal facility.

The effect of diffraction around the earth's curvature is to make possible transmission beyond the line-of-sight. The magnitude of the loss caused by the obstruction increases as either the distance or the frequency is increased and it depends to some extent on the antenna height.¹² The loss resulting from the curvature of the earth is indicated by Fig. 5 as long as neither antenna is higher than the limiting value shown at the top of the chart. This loss is in addition to the transmission loss over plane earth obtained from Fig. 2.

When either antenna is as much as twice as high as the limiting value shown on Fig. 5, this method of correcting for the curvature of the earth indicates a loss that is too great by about 2 db, with the error increasing as the antenna height increases. An alternate method of determining the effect of the earth's curvature is given by Fig. 6. The latter method is approximately correct for any antenna height, but it is theoretically limited in distance to points at or beyond the line-of-sight, assuming that the curved earth is the only obstruction. Fig. 6 gives the loss relative to free-space transmission (and hence is used with Fig. 1) as a func-

tion of three distances: d_1 is the distance to the horizon from the lower antenna, d_2 is the distance to the horizon from the higher antenna, and d_3 is the distance beyond the line-of-sight. In other words, the total distance between antennas, $d = d_1 + d_2 + d_3$. The distance to the horizon over smooth earth is given by:

$$d_{1,2} = \sqrt{2ka h_{1,2}} \quad (8)$$

where $h_{1,2}$ is the appropriate antenna height and ka is the effective earth's radius.

The preceding discussion assumes that the earth is a perfectly smooth sphere and the results are critically dependent on a smooth surface and a uniform atmosphere. The modification in these results caused by the

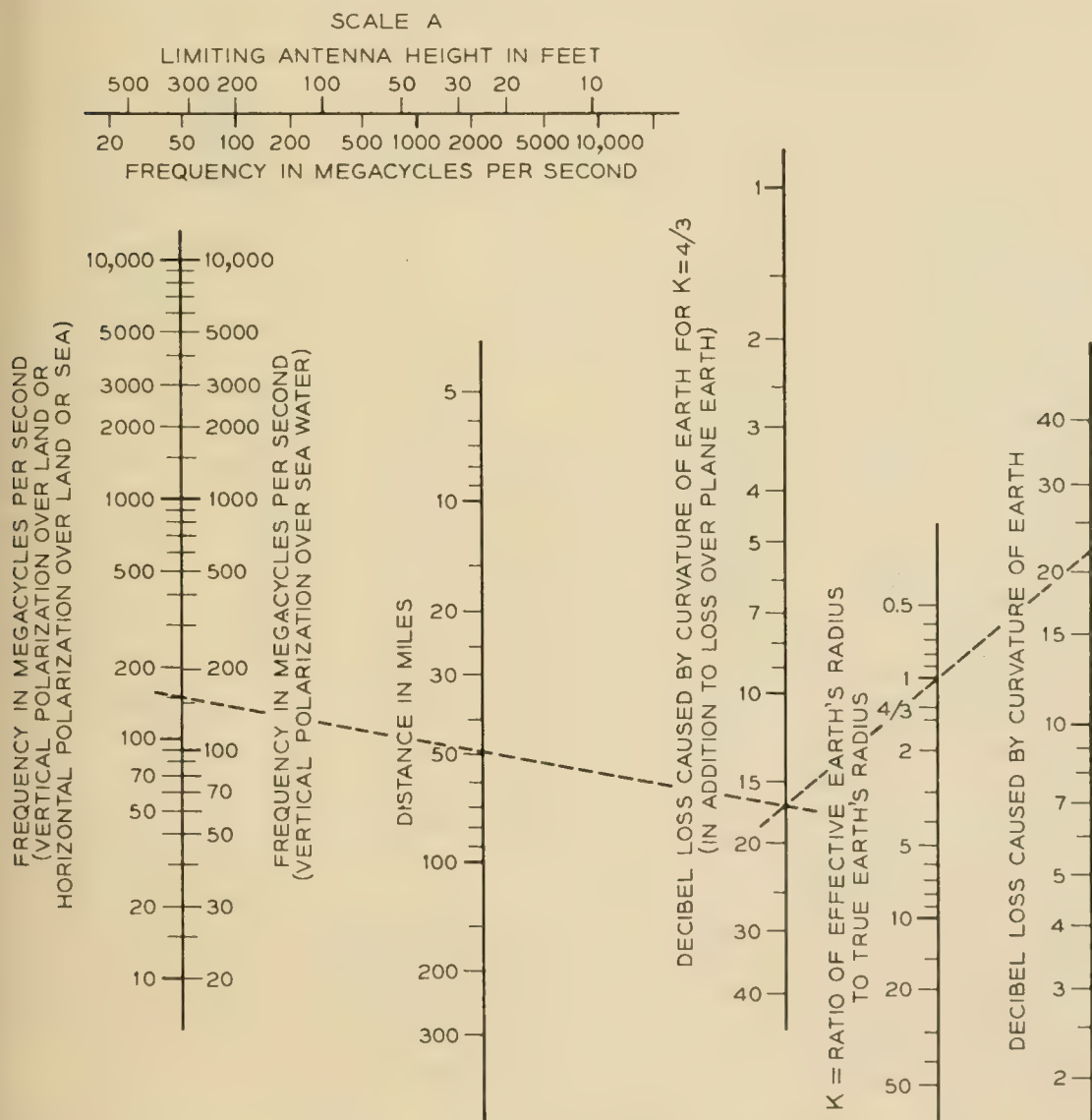


FIG. 5 — Diffraction loss around a perfect sphere.

presence of hills, trees, and buildings is difficult or impossible to compute, but the order of magnitude of these effects may be obtained from a consideration of the other extreme case, which is propagation over a perfectly absorbing knife edge.

The diffraction of plane waves over a knife edge or screen causes a shadow loss whose magnitude is shown on Fig. 7. The height of the obstruction H is measured from the line joining the two antennas to the top of the ridge. It will be noted that the shadow loss approaches 6 db

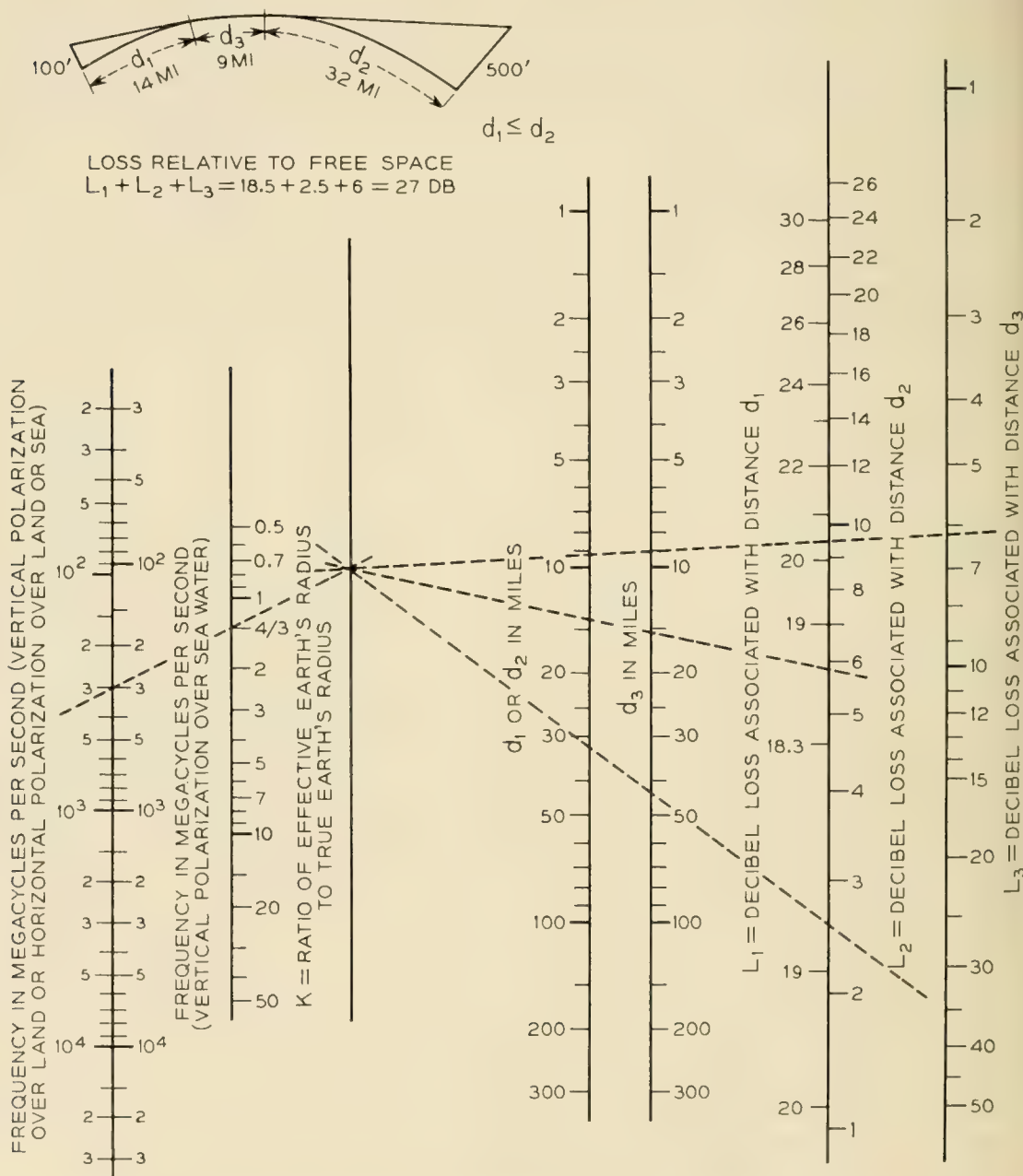


FIG. 6 — Diffraction loss relative to free space transmission at all locations beyond line-of-sight over a smooth sphere.

as H approaches 0 (grazing incidence), and that it increases with increasing positive values of H . When the direct ray clears the obstruction, H is negative, and the shadow loss approaches 0 db in an oscillatory manner as the clearance is increased. In other words, a substantial clearance is required over line-of-sight paths in order to obtain “free-space” transmission. The knife edge diffraction calculation is substantially independent of polarization as long as the distance from the edge is more than a few wavelengths.

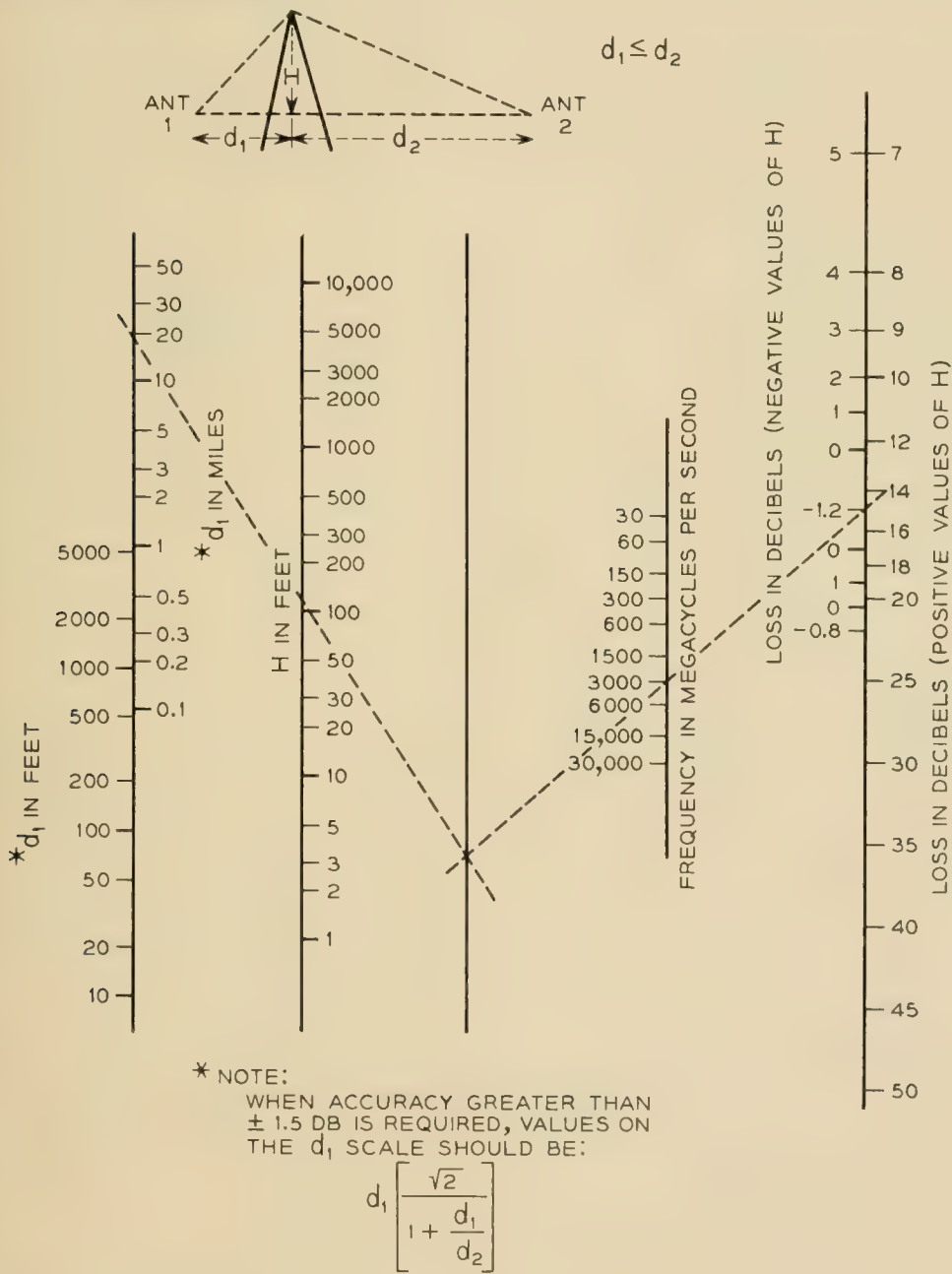


FIG. 7 — Knife-edge diffraction loss relative to free space.

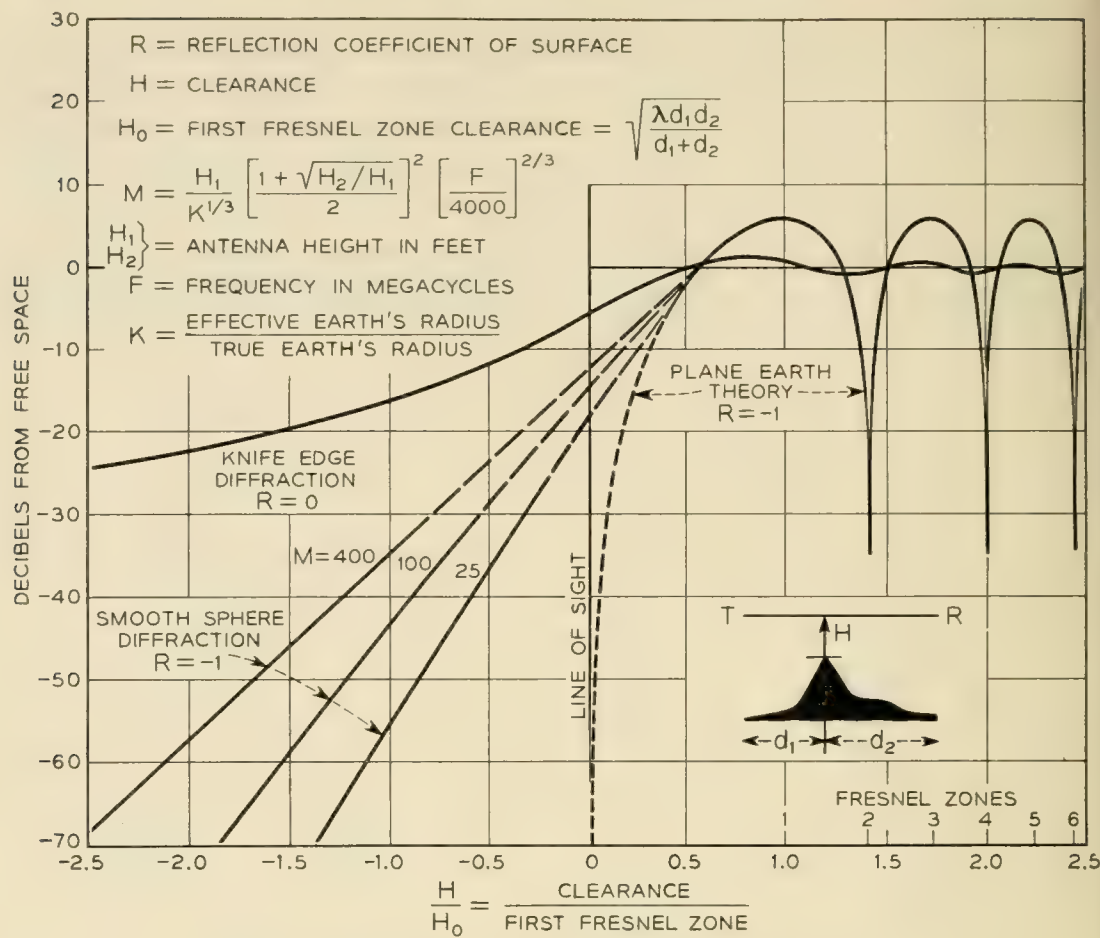


FIG. 8 — Transmission loss versus clearance.

At grazing incidence, the expected loss over a ridge is 6 db (Fig. 7) while over a smooth spherical earth Fig. 6 indicates a loss of about 20 db. More accurate results in the vicinity of the horizon can be obtained by expressing radio transmission in terms of path clearance measured in Fresnel zones as shown in Fig. 8. In this representation the plane earth theory and the ridge diffraction can be represented by single lines; but the smooth sphere theory requires a family of curves with a parameter M that depends primarily on antenna heights and frequency. The big difference in the losses predicted by diffraction around a perfect sphere and by diffraction over a knife edge indicates that diffraction losses depend critically on the assumed type of profile. A suitable solution for the intermediate problem of diffraction over a rough earth has not yet been obtained.

Experimental Data Far Beyond the Horizon

Most of the experimental data at points far beyond the horizon fall in between the theoretical curves for diffraction over a smooth sphere

and for diffraction over a knife edge obstruction. Various theories have been advanced to explain these effects but none has been reduced to a simple form for every day use.¹³ The explanation most commonly accepted is that energy is reflected or scattered from turbulent air masses in the volume of air that is enclosed by the intersection of the beamwidths of the transmitting and receiving antennas.¹⁴

The variation in the long term median signals with distance has been derived from experimental results and is shown in Fig. 9 for two frequencies.¹⁵ The ordinate is in db below the signal that would have been expected at the same distance in free space with the same power and the same antennas. The strongest signals are obtained by pointing the antennas at the horizon along the great circle route. The values shown on Fig. 9 are essentially annual averages taken from a large number of paths, and substantial variations are to be expected with terrain, climate, and season as well as from day to day fading.

Antenna sites with sufficient clearance so that the horizon is several miles away will, on the average, provide a higher median signal (less loss) than shown on Fig. 9. Conversely, sites for which the antenna must be pointed upward to clear the horizon will ordinarily result in appreciably more loss than shown on Fig. 9. In many cases the effects of path length and angles to the horizon can be combined by plotting the experimental results as a function of the angle between the lines drawn tangent to the horizon from the transmitting and receiving sites.¹⁶

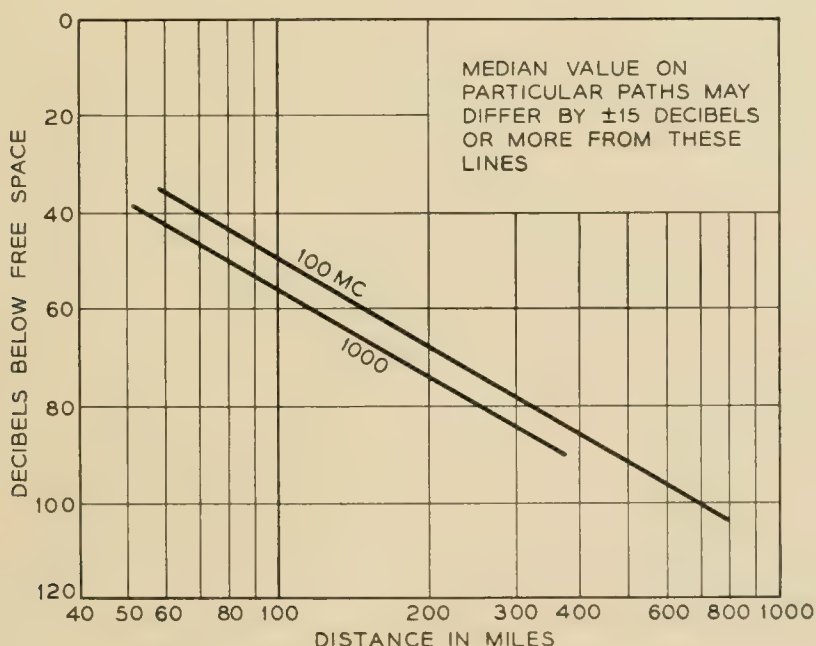


FIG. 9 — Beyond-horizon transmission — median signal level versus distance.

When the path profile consists of a single sharp obstruction that can be seen from both terminals, the signal level may approach the value predicted by the knife edge diffraction theory.¹⁷ While several interesting and unusual cases have been recorded, the knife edge or "obstacle gain" theory is not applicable to the typical but only to the exceptional paths.

As in the case of line-of-sight transmission the fading of radio signals beyond the horizon can be divided into fast fading and slow fading. The fast fading is caused by multipath transmission in the atmosphere, and for a given size antenna, the rate of fading increases as either the frequency or the distance is increased. This type of fading is much faster than the maximum fast fading observed on line of sight paths, but the two are similar in principle. The magnitude of the fades is described by the Rayleigh distribution.

Slow fading means variations in average signal level over a period of hours or days and it is greater on beyond horizon paths than on line-of-sight paths. This type of fading is almost independent of frequency and seems to be associated with changes in the average refraction of the atmosphere. At distances of 150 to 200 miles the variations in hourly median value around the annual median seem to follow a normal probability law in db with a standard deviation of about 8 db. Typical fading distributions are shown on Fig. 10.

The median signal levels are higher in warm humid climates than in cold dry climates and seasonal variations of as much as ± 10 db or more from the annual median have been observed.¹⁸

Since the scattered signals arrive with considerable phase irregularities in the plane of the receiving antenna, narrow-beamed (high gain) antennas do not yield power outputs proportional to their theoretical area gains. This effect has sometimes been called loss in antenna gain, but it is a propagation effect and not an antenna effect. On 150 to 200 miles this loss in received power may amount to one or two db for a 40 db gain antenna, and perhaps six to eight db for a 50 db antenna. These extra losses vary with time but the variations seem to be uncorrelated with the actual signal level.

The bandwidth that can be used on a single radio carrier is frequently limited by the selective fading caused by multipath or echo effects. Echoes are not troublesome as long as the echo time delays are very short compared with one cycle of the highest baseband frequency. The probability of long delayed echoes can be reduced (and the rate of fast fading can be decreased) by the use of narrow beam antennas both within and beyond the horizon.^{19, 20} Useful bandwidths of several mega-

cycles appear to be feasible with the antennas that are needed to provide adequate signal-to-noise margins. Successful tests of television and of multichannel telephone transmission have been reported on a 188-mile path at 5,000 mc.²¹

The effects of fast fading can be reduced substantially by the use of either frequency or space diversity. The frequency or space separation required for diversity varies with time and with the degree of correlation that can be tolerated. A horizontal (or vertical) separation of about 100 wavelengths is ordinarily adequate for space diversity on 100- to 200-mile paths. The corresponding figure for the required frequency separation for adequate diversity seems likely to be more than 20 mc.

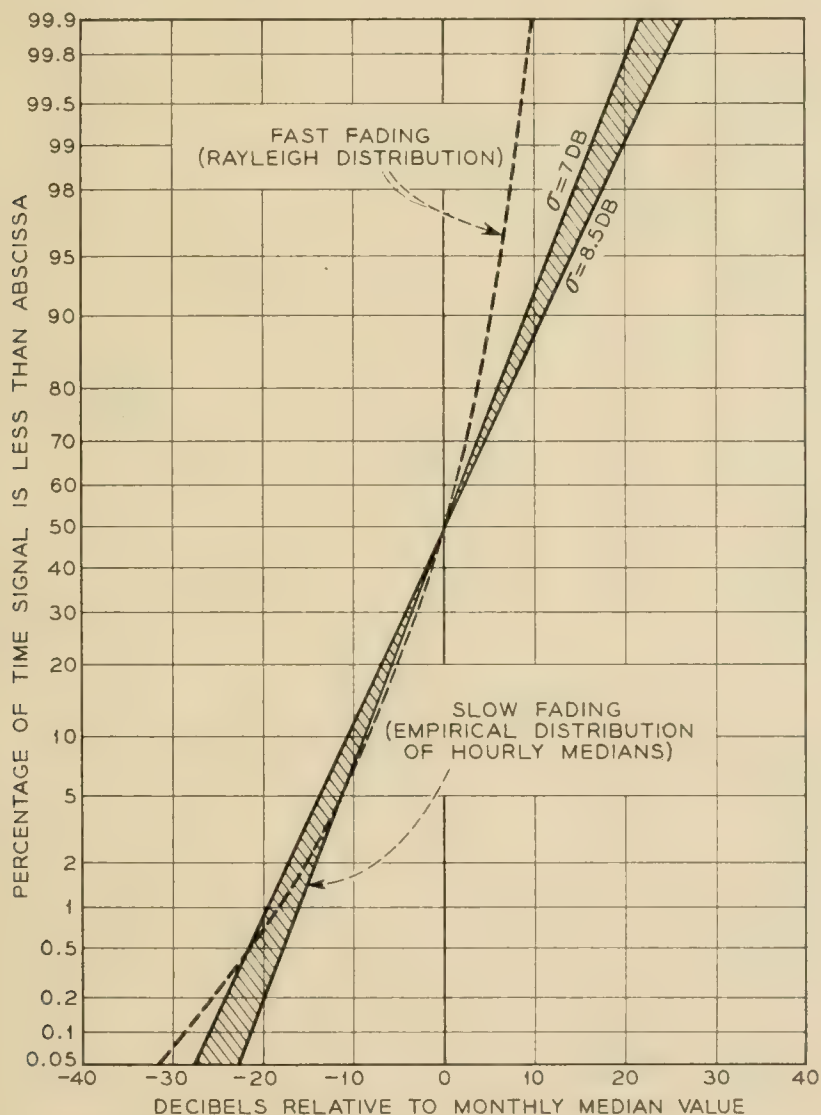


FIG. 10 — Typical fading characteristics at points far beyond the horizon.

Effects of Nearby Hills — Particularly on Short Paths

The experimental results on the effects of hills indicate that the shadow losses increase with the frequency and with the roughness of the terrain.²²

An empirical summary of the available data is shown on Fig. 11. The roughness of the terrain is represented by the height H shown on the profile at the top of the chart. This height is the difference in elevation between the bottom of the valley and the elevation necessary to obtain line of sight from the transmitting antenna. The right hand scale in Fig. 11 indicates the additional loss above that expected over plane earth. Both the median loss and the difference between the median and the 10 per cent values are shown. For example, with variations in terrain of 500 feet, the estimated median shadow loss at 450 mc is about 20 db and the

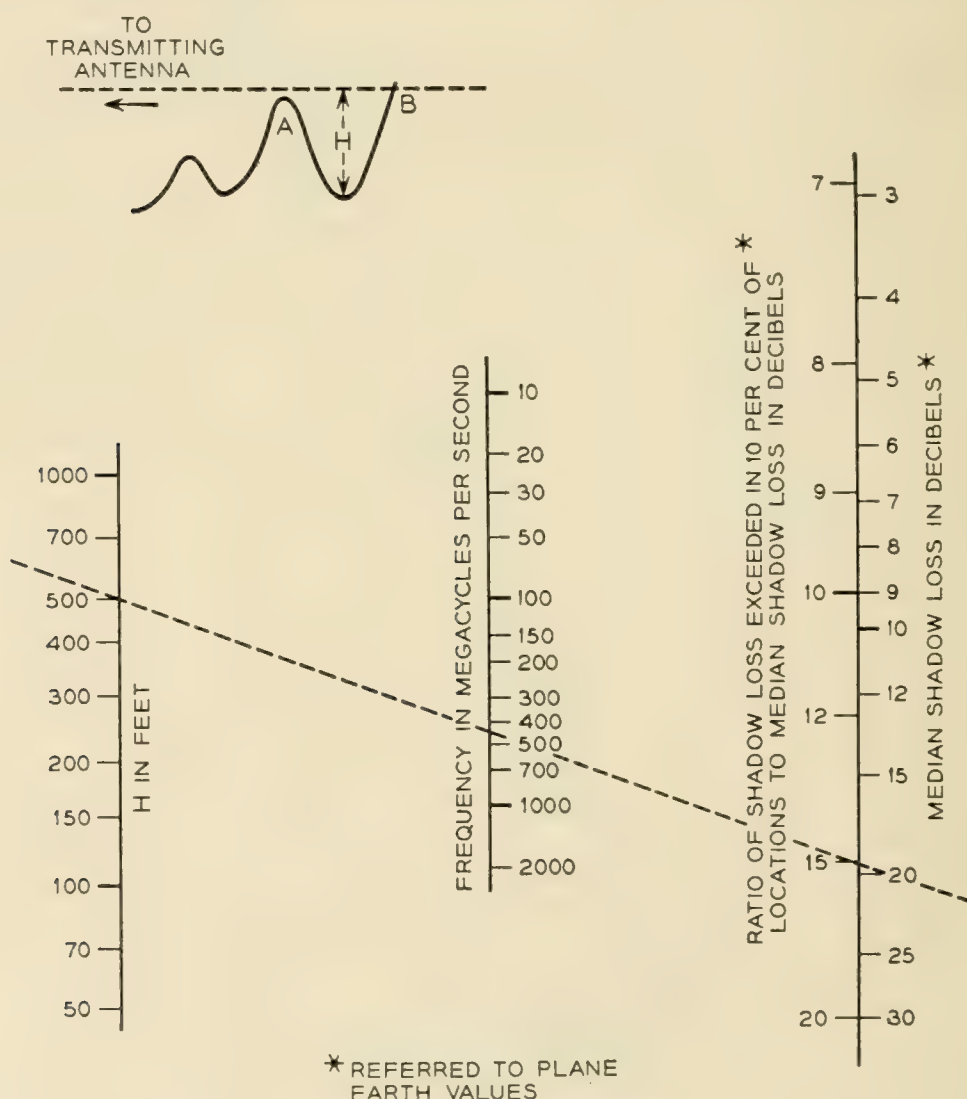


FIG. 11 — Estimated distribution of shadow losses.

shadow loss exceeded in only 10 per cent of the possible locations between points *A* and *B* is about $20 + 15 = 35$ db. It will be recognized that this analysis is based on large-scale variations in field intensity, and does not include the standing wave effects which sometimes cause the field intensity to vary considerably within a few feet.

Effects of Buildings and Trees

The shadow losses resulting from buildings and trees follow somewhat different laws from those caused by hills. Buildings may be more transparent to radio waves than the solid earth, and there is ordinarily much more back scatter in the city than in the open country. Both of these factors tend to reduce the shadow losses caused by the buildings but, on the other hand, the angles of diffraction over or around the buildings are usually greater than for natural terrain. In other words, the artificial canyons caused by buildings are considerably narrower than natural valleys, and this factor tends to increase the loss resulting from the presence of buildings. The available quantitative data on the effects of buildings are confined primarily to New York City. These data indicate that in the range of 40 to 450 mc there is no significant change with frequency, or at least the variation with frequency is somewhat less than that noted in the case of hills.²³ The median field intensity at street level for random locations in Manhattan (New York City) is about 25 db below the corresponding plane earth value. The corresponding values for the 10 per cent and 90 per cent points are about 15 and 35 db, respectively.

Typical values of attenuation through a brick wall, are from 2 to 5 db at 30 mc and 10 to 40 db at 3,000 mc, depending on whether the wall is dry or wet. Consequently most buildings are rather opaque at frequencies of the order of thousands of megacycles.

When an antenna is surrounded by moderately thick trees and below tree-top level, the average loss at 30 mc resulting from the trees is usually 2 or 3 db for vertical polarization and is negligible with horizontal polarization. However, large and rapid variations in the received field intensity may exist within a small area, resulting from the standing-wave pattern set up by reflections from trees located at a distance of several wavelengths from the antenna. Consequently, several near-by locations should be investigated for best results. At 100 mc the average loss from surrounding trees may be 5 to 10 db for vertical polarization and 2 or 3 db for horizontal polarization. The tree losses continue to increase as the frequency increases, and above 300 to 500 mc they tend to be independent of the type of polarization. Above 1,000 mc, trees that are thick

enough to block vision are roughly equivalent to a solid obstruction of the same over-all size.

IV. MEDIUM AND LOW FREQUENCY GROUND WAVE TRANSMISSION

Wherever the antenna heights are small compared with the wavelength, the received field intensity is ordinarily stronger with vertical polarization than with horizontal and is stronger over sea water than over poor soil. In these cases the "surface wave" term in (3) cannot be neglected. This use of the term "surface wave" follows Norton's usage and is not equivalent to the Sommerfeld or Zenneck "surface waves."

The parameter A is the plane earth attenuation factor for antennas at ground level. It depends upon the frequency, ground constants, and type of polarization. It is never greater than unity and decreases with increasing distance and frequency, as indicated by the following approximate equation:^{24, 25}

$$A \approx \frac{-1}{1 + j \frac{2\pi d}{\lambda} (\sin \theta + z)^2} \quad (9)$$

where

$$z = \frac{\sqrt{\epsilon_0 - \cos^2 \theta}}{\epsilon_0} \text{ for vertical polarization}$$

$$z = \sqrt{\epsilon_0 - \cos^2 \theta} \text{ for horizontal polarization}$$

$$\epsilon_0 = \epsilon - j60\sigma\lambda$$

$$\theta = \text{angle between reflected ray and the ground}$$

$$= 0 \text{ for antennas at ground level}$$

$$\epsilon = \text{dielectric constant of the ground relative to unity in free space}$$

$$\sigma = \text{conductivity of the ground in mhos per meter}$$

$$\lambda = \text{wavelength in meters}$$

In terms of these same parameters the reflection coefficient of the ground is given by²⁶

$$R = \frac{\sin \theta - z}{\sin \theta + z} \quad (10)$$

When $\theta \ll |z|$ the reflection coefficient approaches -1 ; when $\theta \gg |z|$

(which can happen only with vertical polarization) the reflection coefficient approaches $+1$. The angle for which the reflection coefficient is a minimum is called the pseudo-Brewster angle and it occurs for $\sin \theta = |z|$.

For antennas approaching ground level the first two terms in (3) cancel each other (h_1 and h_2 approach zero and R approaches -1) and the magnitude of the third term becomes

$$|(1 - R)A| \approx \frac{2}{\frac{2\pi d}{\lambda} z^2} = \frac{4\pi h_0^2}{\lambda d} \quad (11)$$

where $h_0 =$ minimum effective antenna height shown in Fig. 3

$$= \left| \frac{\lambda}{2\pi z} \right|$$

The surface wave term arises because the earth is not a perfect reflector. Some energy is transmitted into the ground and sets up ground currents, which are distorted relative to what would have been the case in an ideal perfectly reflecting surface. The surface wave is defined as the vertical electric field for vertical polarization, or the horizontal electric field for horizontal polarization, that is associated with the extra components of the ground currents caused by lack of perfect reflection. Another component of the electric field associated with the ground currents is in the direction of propagation. It accounts for the success of the wave antenna at lower frequencies, but it is always smaller in magnitude than the surface wave as defined above. The components of the electric vector in three mutually perpendicular co-ordinates are given by Norton.²⁷

In addition to the effect of the earth on the propagation of radio waves, the presence of the ground may also affect the impedance of low antennas and thereby may have an effect on the generation and reception of radio waves.²⁸ As the antenna height varies, the impedance oscillates around the free space value, but the variations in impedance are usually unimportant as long as the center of the antenna is more than a quarter-wavelength above the ground. For vertical grounded antennas (such as are used in standard AM broadcasting) the impedance is doubled and the net effect is that the maximum field intensity is 3 db above the free space value instead of 6 db as indicated in (4) for elevated antennas.

Typical values of the field intensity to be expected from a grounded quarter-wave vertical antenna are shown in Fig. 12 for transmission over poor soil and in Fig. 13 for transmission over sea water. These charts in-

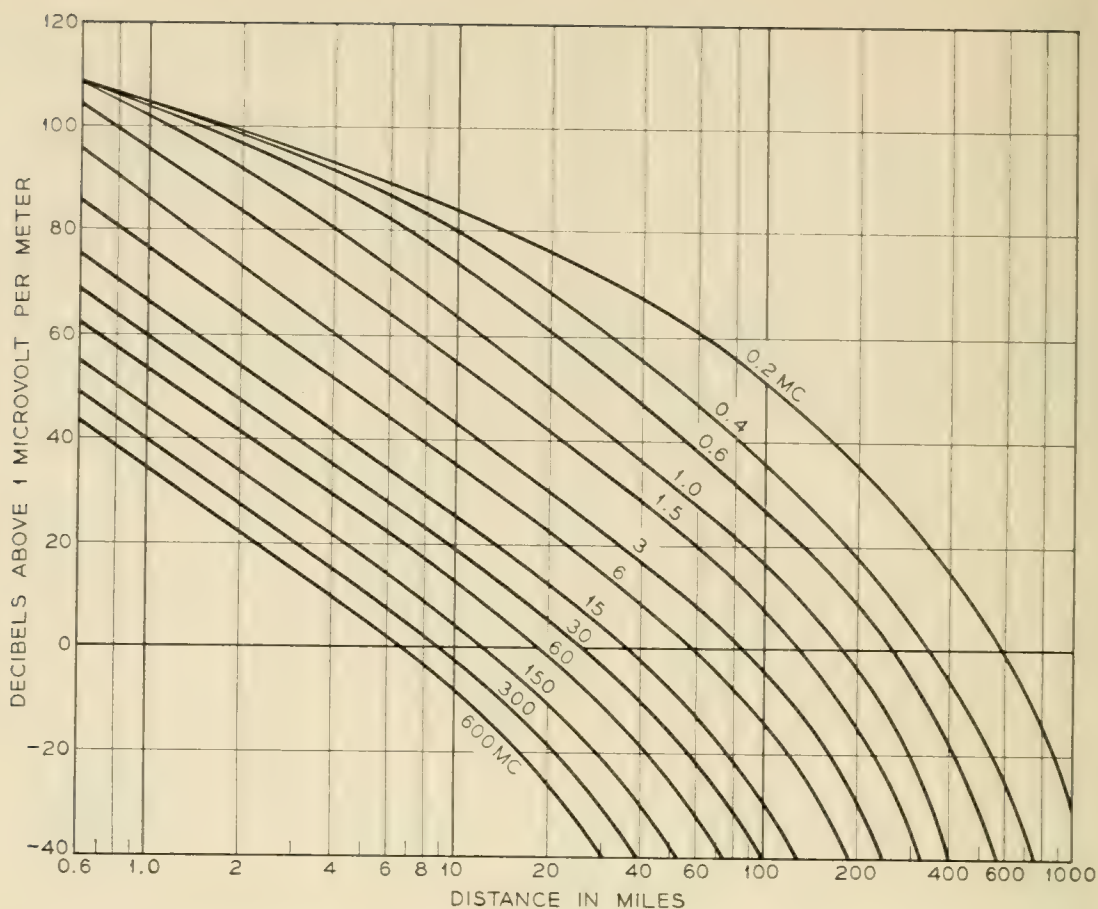


FIG. 12 — Field intensity for vertical polarization over poor soil for 1-kw radiated power from a grounded whip antenna.

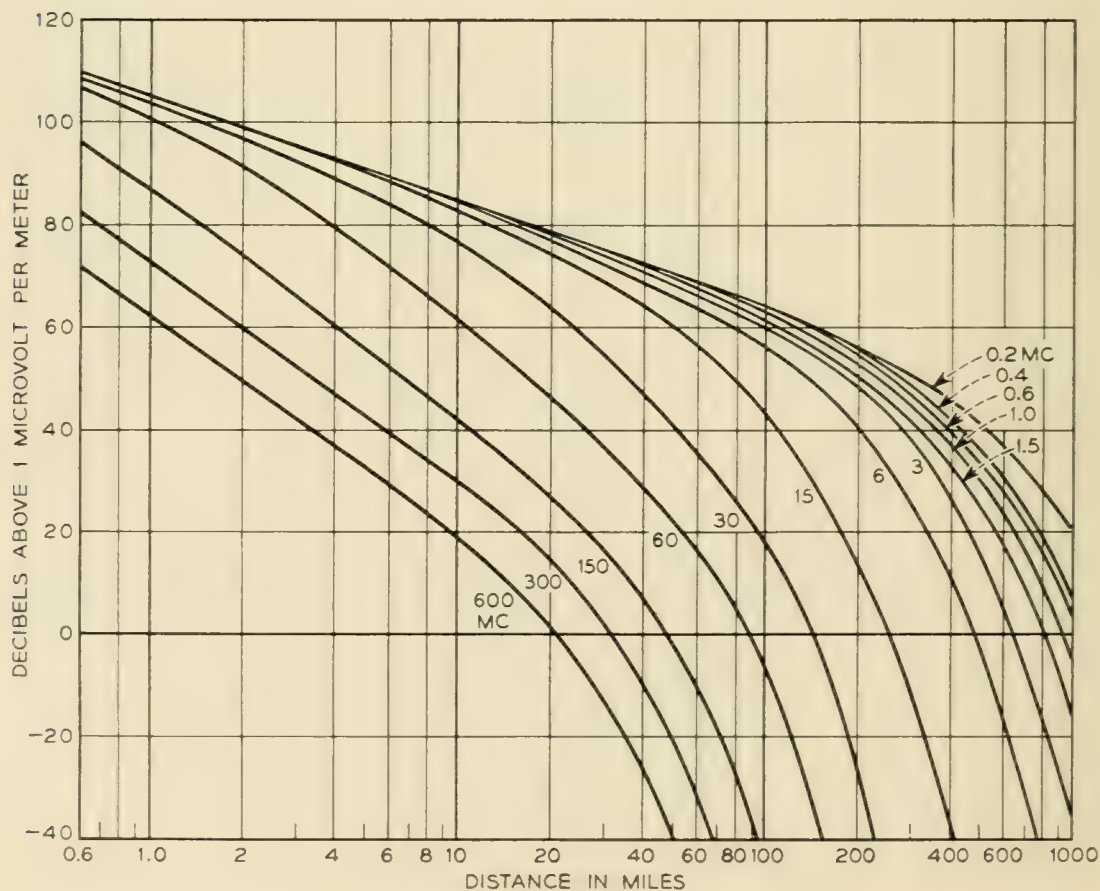


FIG. 13 — Field intensity for vertical polarization over sea water for 1-kw radiated power from a grounded whip antenna.

clude the effect of diffraction and average refraction around a smooth spherical earth as discussed in Section III, but do not include the ionospheric effects described in the next Section. The increase in signal obtained by raising either antenna height is shown in Fig. 14 for poor soil and Fig. 15 for sea water.

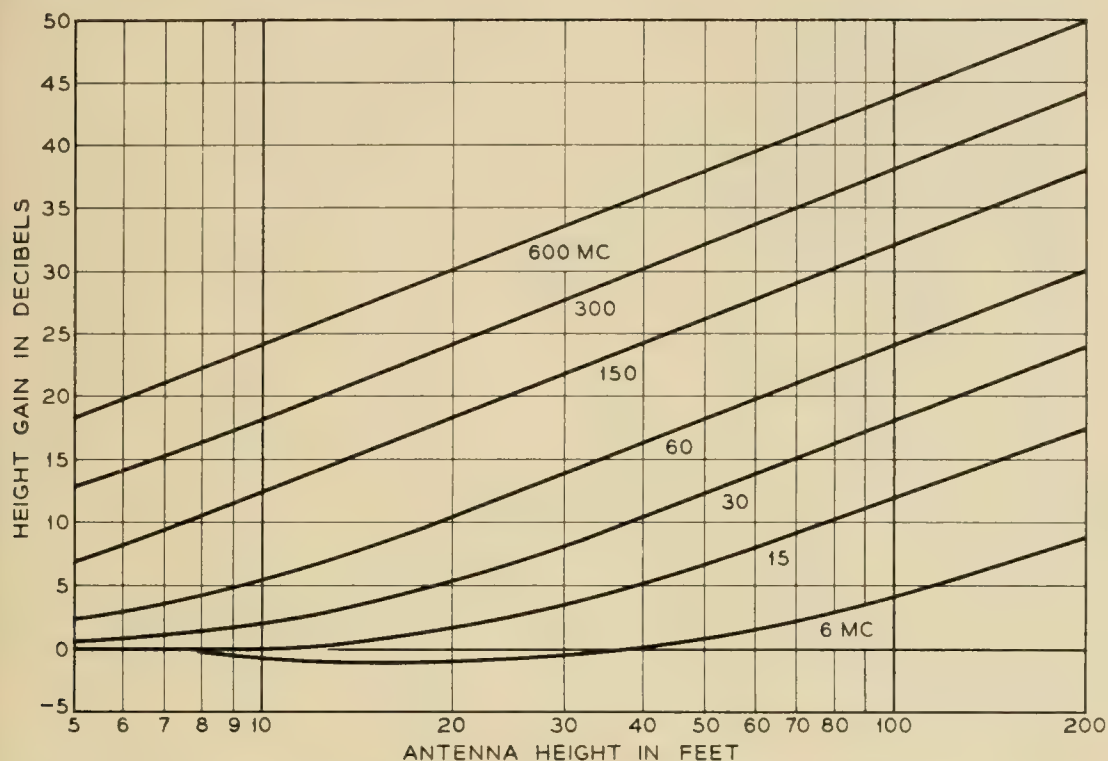


FIG. 14 — Antenna height gain factor for vertical polarization over poor soil.

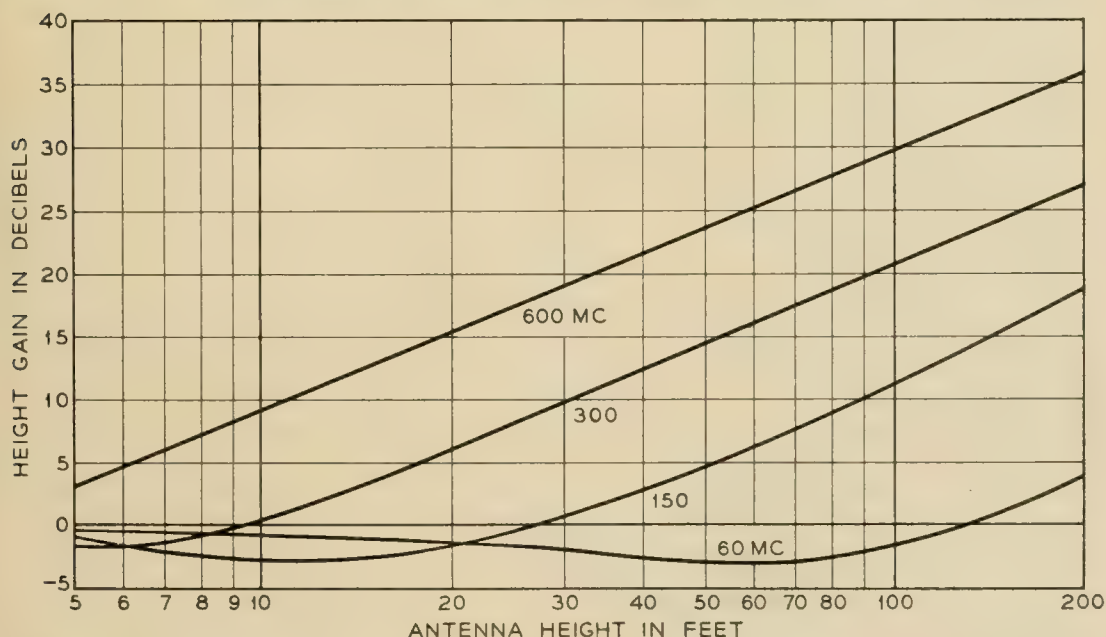


FIG. 15 — Antenna height gain factor for vertical polarization over sea water.

V. IONOSPHERIC TRANSMISSION

In addition to the tropospheric or ground wave transmission discussed in the preceding sections, useful radio energy at frequencies below about 25 to 100 mc may be returned to the earth by reflection from the ionosphere, which consists of several ionized layers located 50 to 200 miles above the earth. The relatively high density of ions and free electrons in this region provides an effective index of refraction of less than one, and the resulting transmission path is similar to that in the well known optical phenomenon of total internal reflection. The mechanism is generally spoken of as reflection from certain virtual heights.²⁹ Polarization is not maintained in ionospheric transmission and the choice depends on the antenna design that is most efficient at the desired elevation angles.

Regular Ionospheric Transmission

The ionosphere consists of three or more distinct layers. This does not mean that the space between layers is free of ionization but rather that the curve of ion density versus height has several distinct peaks. The E , F_1 , and F_2 layers are present during the daytime but the F_1 and F_2 combine to form a single layer at night. A lower layer called the D layer is also present during the day, but its principal effect is to absorb rather than reflect.

Information about the nature of the ionosphere has been obtained by transmitting pulsed radio signals directly overhead and by recording the signal intensity and the time delay of the echoes returned from these layers. At night all frequencies below the critical frequency f_c are returned to earth with an average signal intensity that is about 3 to 6 db below the free space signal that would be expected for the round trip distance. At frequencies higher than the critical frequency the signal intensity is very weak or undetectable. Typical values of the critical frequency for Washington, D. C., are shown in Fig. 16.

During the daytime, the critical frequency is increased 2 to 3 times over the corresponding nighttime value. This apparent increase in the useful frequency range for ionospheric transmission is largely offset by the heavy daytime absorption which reaches a maximum in the 1 to 2-mc range. This absorption is caused by interaction between the free electrons and the earth's magnetic field. The absence of appreciable absorption at night indicates that most of the free electrons disappear when the sun goes down. Charged particles traveling in a magnetic field have a resonant or gyro-magnetic frequency, and for electrons in the earth's magnetic field, of about 0.5 gauss, this resonance occurs at about 1.4

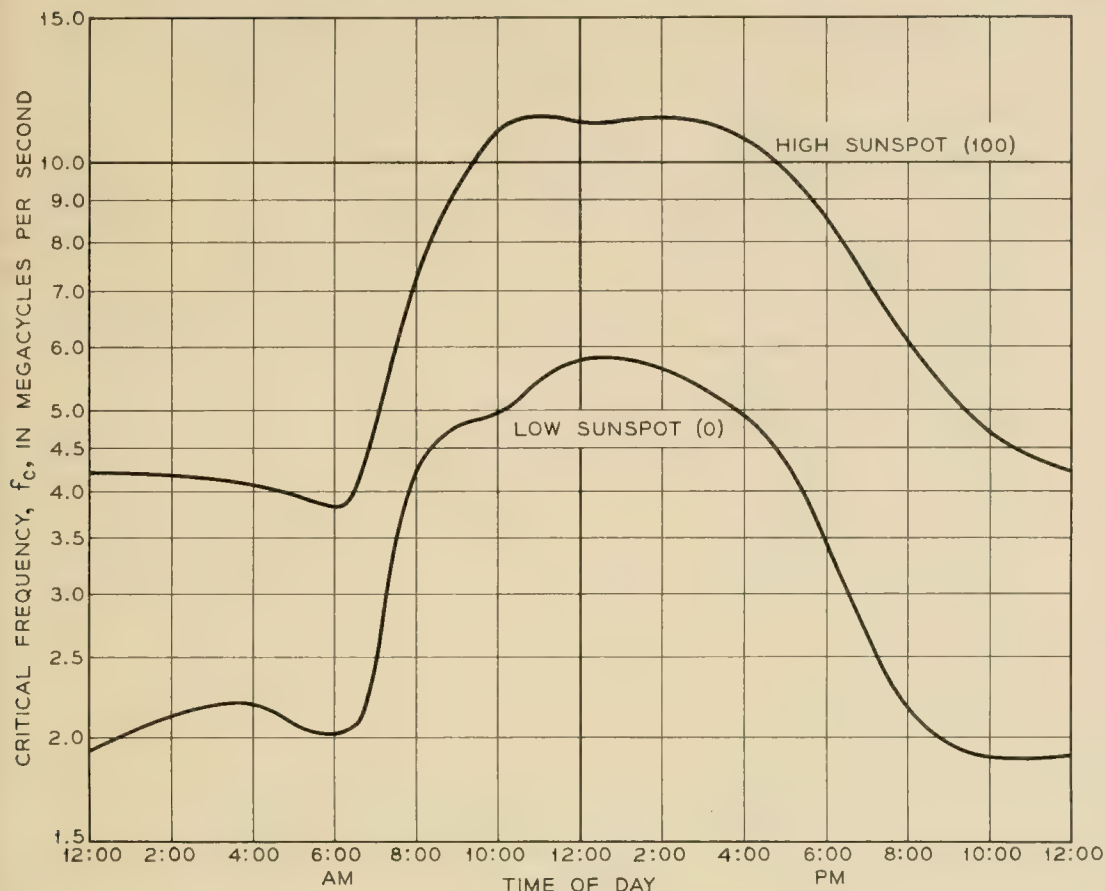


FIG. 16 — Typical diurnal variation of critical frequency for January at latitude 40 degrees.

mc. The magnitude of the absorption varies with the angle of the sun above the horizon and is a maximum about noon. The approximate midday absorption is shown on Fig. 17 in terms of db per 100 miles of path length. (On short paths this length is the actual path traveled, not the distance along the earth's surface.)

Long distance transmission requires that the signal be reflected from the ionosphere at a small angle instead of the perpendicular incidence used in obtaining the critical frequency. For angles other than directly overhead an assumption which seems to be borne out in practice is that the highest frequency for which essentially free space transmission is obtained is $f_c/\sin \alpha$, where α is the angle between the radio ray and ionospheric layer. This limiting frequency is greater than the critical frequency and is called the maximum usable frequency which is usually abbreviated muf. The curved geometry limits the distance that can be obtained with one-hop transmission to about 2,500 miles and the muf at the longer distances does not exceed 3 to 3.5 times the critical frequency.

The difference between day and night effects means that most sky-

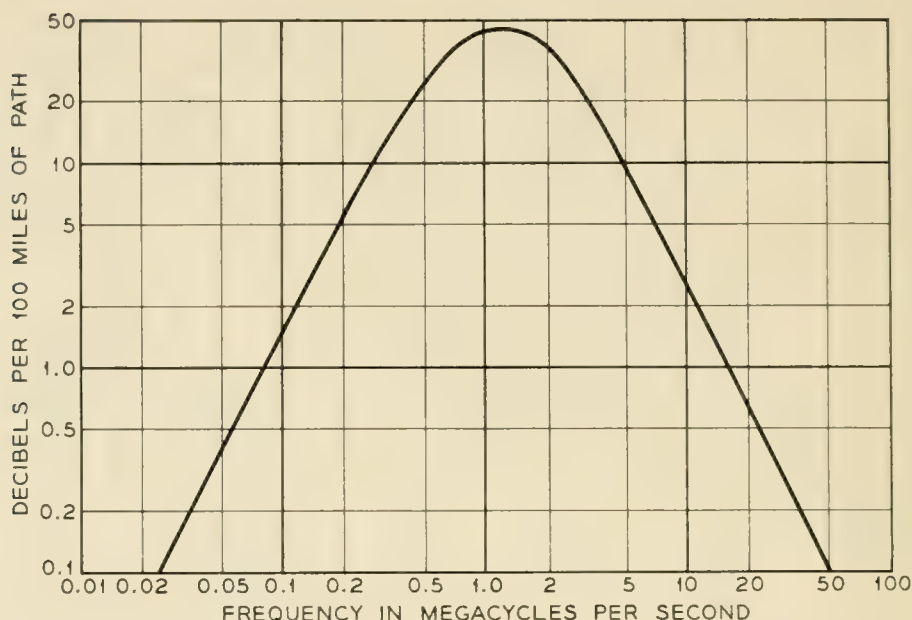


FIG. 17 — Typical values of midday ionospheric absorption.

wave paths require at least two frequencies. A relatively low frequency is needed to get under the nighttime muf and a higher frequency is needed that is below the daytime muf but above the region of high absorption. This lower limit depends on the available signal-to-noise margin and is commonly called the lowest useful high frequency.

Frequencies most suitable for transmission of 1000 miles or more will ordinarily not be reflected at the high angles needed for much shorter distances. As a result the range of skywave transmission ordinarily does not overlap the range of groundwave transmission, and the intermediate region is called the skip zone because the signal is too weak to be useful. At frequencies of a few megacycles the groundwave and skywave ranges may overlap with the result that severe fading occurs when the two signals are comparable in amplitude.

In addition to the diurnal variations in frequency and in absorption there are systematic changes with season, latitude, and with the nominally eleven-year sunspot cycle. Random changes in the critical frequency of about ± 15 per cent from the monthly median value are also to be expected from day to day.

The *F* layer is the principal contributor to transmission beyond 1,000 to 1,500 miles and typical values of the maximum usable frequency can be summarized as follows: The median nighttime critical frequency for *F* layer transmission at the latitude at Washington, D. C., is about 2 mc in the month of June during a period of low sunspot activity. All frequencies below about 2 mc are strongly reflected to earth while the higher

frequencies are either greatly attenuated or are lost in outer space. The approximate maximum usable frequency for other conditions is greater than 2 mc by the ratios shown in Table I.

TABLE I

Variation With	Multiplying factor
(1) Time of Day	
Midnight	1 (Reference)
Early Afternoon—June	2
December	3
(2) Path Length	
Less than 200 Miles	1 (Reference)
Approx. 1000 Miles	2
More than 2500 Miles	3.5
(3) Sunspot Cycle	
Minimum	1 (Reference)
For one year in five—June	1.5
December	2
For one year in fifty—June	2
December	3

When all of the above variations add “in phase,” transmission for distances of 2,500 miles or more is possible at frequencies up to 40 to 60 mc. For example, using the table, 2,500-mile transmission on an early December afternoon in one year out of five can be expected on a frequency of about 42 mc, which is $3 \times 3.5 \times 2 = 21$ times the reference critical frequency of 2 mc. Peaks of the sunspot cycle occurred in 1937 and in 1947–1948 so another peak is expected in 1958–1959.

The maximum usable frequency also varies with the geomagnetic latitude but, as a first approximation, the above values are typical of continental U. S. Forecasts of the muf to be expected throughout the world are issued monthly by the National Bureau of Standards.^{30, 31} These estimates include the diurnal, seasonal, and sunspot effects.

Another type of absorption, over and above the usual daytime absorption, occurs both day and night on transmission paths that travel through the auroral zone. The auroral zones are centered on the north and south magnetic poles at about the same distance as the Arctic Circle is from the geographical north pole. During periods of magnetic storms these auroral zones expand over an area much larger than normal and thereby disrupt communication by introducing unexpected absorption. These conditions of poor transmission can last for hours and sometimes even for days. These periods of increased absorption are more common in the polar regions than in the temperate zones or the tropics because of the proximity of the auroral zone and are frequently called HF “black-outs.” During a “blackout,” the signal level is decreased considerably

but the signal does not drop out completely. It appears possible that the outage time normally associated with HF transmission could be greatly reduced by the use of transmitter power and antenna size comparable to that needed in the ionospheric scatter method described below.

In addition to the auroral zone absorption, there are shorter periods of severe absorption over the entire hemisphere facing the sun. These erratic and unpredictable effects which seem to be associated with eruptions on the sun are called sudden ionospheric disturbances (SID's) or the Dellinger effect.

The preceding information is based primarily on F layer transmission. The E layer is located closer to the earth than the F layer and the maximum transmission distance for a single reflection is about 1,200 miles.

Reflections from the E layer sometimes occur at frequencies above about 20 mc but are erratic in both time and space. This phenomenon has been explained by assuming that the E layer contains clouds of ionization that are variable in size, density, and location. The maximum frequency returned to earth may at times be as high as 70 or 80 mc.³² The high values are more likely to occur during the summer, and during the minimum of the sunspot cycle.

Rapid multipath fading exists on ionospheric circuits and is superimposed on the longer term variations discussed above. The amplitude of the fast fading follows the Rayleigh distribution and echo delays up to several milliseconds are observed. These delays are 10^4 to 10^5 times as long as for tropospheric transmission. As a result of these relatively long delays uncorrelated selective fading can occur within a few hundred cycles. This produces the distortion on voice circuits that is characteristic of "short wave" transmission.

Ionospheric Scatter

The maximum usable frequency used in conventional skywave transmission is defined as the highest frequency returned to earth for which the average transmission is within a few db of free space. As the frequency increases above the muf the signal level decreases rapidly but does not drop out completely. Although the signal level is low, reliable transmission can be obtained at frequencies up to 50 mc or higher and to distances up to at least 1,200 to 1,500 miles.³³ In this case the signal is 80 to 100 db below the free space value and its satisfactory use requires much higher power and larger antennas than are ordinarily used in ionospheric transmission. The approximate variation in median signal level with frequency is shown in Fig. 18.

Ionospheric scatter is apparently the result of reflections from many

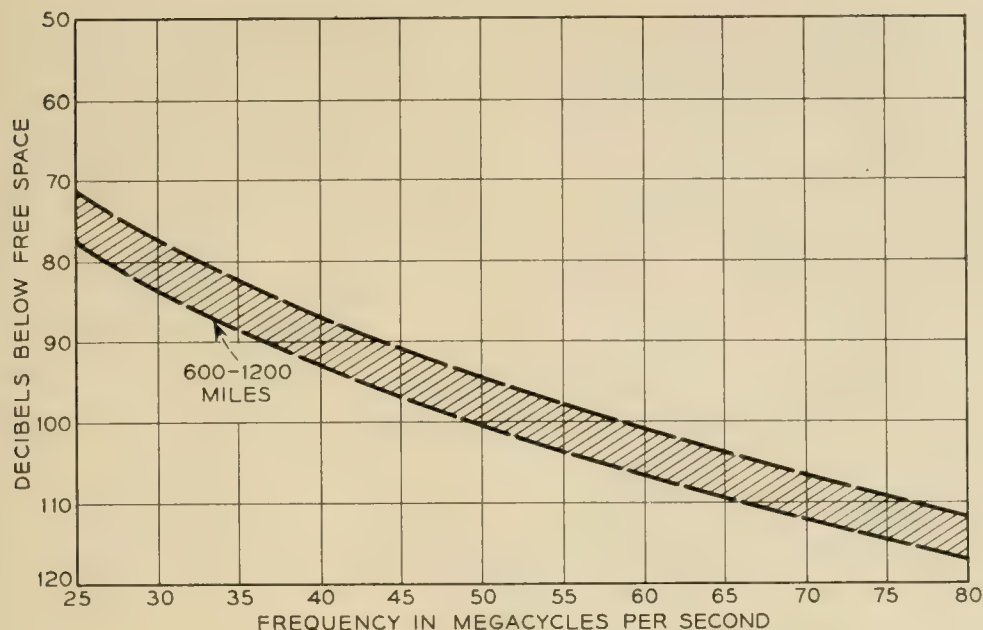


FIG. 18 — Median signal levels for ionospheric scatter transmission.

patches of ionization in the *E* layer. It is suspected that meteors are important in establishing and in maintaining this ionization but this has not been clearly determined.

In common with other types of transmission, the fast fading follows a Rayleigh distribution. The distribution of hourly median values relative to the long term median (after the high signals resulting from sporadic *E* transmission have been removed) is approximately a normal probability law with a standard deviation of about 6 to 8 db.

Ionospheric scatter transmission is suitable for several telegraph channels but the useful bandwidth is limited by the severe selective fading that is characteristic of all ionospheric transmission.

VI. NOISE LEVELS

The usefulness of a radio signal is limited by the “noise” in the receiver. This noise may be either unwanted external interference or the first circuit noise in the receiver itself.

Atmospheric static is ordinarily controlling at frequencies below a few megacycles while set noise is the primary limitation at frequencies above 200 to 500 mc. In the 10- to 200-mc band the controlling factor depends on the location, time of day, etc. and may be either atmospheric static, man made noise, cosmic noise, or set noise.

The theoretical minimum circuit noise caused by the thermal agitation of the electrons at usual atmospheric temperatures is 204 db below one

watt per cycle of bandwidth; that is, the thermal noise power, in dbw is $-204 + 10 \log (\text{bandwidth})$. The first circuit or set noise is usually higher than the theoretical minimum by a factor known as the noise figure. For example, the set noise in a receiver with a 6-kc noise bandwidth and an 8-db noise figure is 158 db below 1 watt, which is equivalent to 0.12 microvolts across 100 ohms. Variations in thermal noise and set noise follow the Rayleigh distribution, but the quantitative reference is usually the rms value (63.2 per cent point), which is 1.6 db higher than the median value shown on Figs. 4 and 10. Momentary thermal noise peaks more than 10 to 12 db above the median value occur for a small percentage of the time.

Atmospheric static is caused by lightning and other natural electrical disturbances, and is propagated over the earth by ionospheric transmission. Static levels are generally stronger at night than in the daytime. Atmospheric static is more noticeable in the warm tropical areas where the storms are most frequent than it is in the colder northern regions which are far removed from the lightning storms.

Typical average values of noise in a 6-kc band are shown on Fig. 19. The atmospheric static data are rough yearly averages for a latitude of 40° . Typical summer averages are a few db higher than the value on Fig. 19 and the corresponding winter values are a few db lower. The average noise levels in the tropics may be as much as 15 db higher than

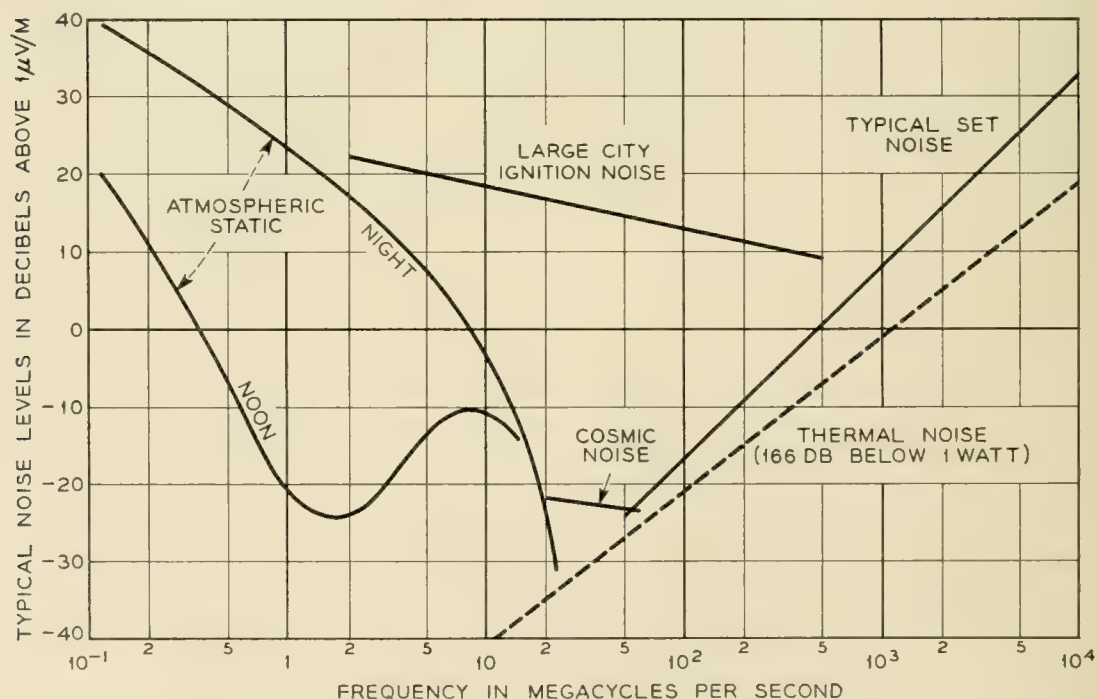


FIG. 19 — Typical average noise level in a 6-kc band.

for latitudes of 40° while in the Arctic region the noise may be 15 to 25 db lower. The corresponding values for other bandwidths can be obtained by adding 10 db for each 10-fold increase in bandwidth. More complete estimates of atmospheric noise on a world wide basis are given in the National Bureau of Standards Bulletin 462.²⁹ These noise data are based on measurements with a time constant of 100 to 200 milliseconds. Noise peaks, as measured on a cathode ray tube, may be considerably higher.

The man made noise shown on Fig. 19 is caused primarily by operation of electric switches, ignition noise, etc., and may be a controlling factor at frequencies below 200 to 400 mc. Since radio transmission in this frequency range is primarily tropospheric (ground wave), man made noise can be relatively unimportant beyond 10 to 20 miles from the source. In rural areas, the controlling factor can be either set noise or cosmic noise.

Cosmic and solar noise is a thermal type interference of extra-terrestrial origin.³⁴ Its practical importance as a limitation on communication circuits seems to be in the 20- to 80-mc range. Cosmic noise has been found at much higher frequencies but its magnitude is not significantly above set noise. On the other hand, noise from the sun increases as the frequency increases and may become the controlling noise source when high gain antennas are used. The rapidly expanding science of radio astronomy is investigating the variations in both time and frequency of these extra-terrestrial sources of radio energy.

REFERENCES

1. H. T. Friis, A Note on a Simple Transmission Formula Proc. I.R.E., **34**, pp. 254-256, May, 1946.
2. K. Bullington, Radio Propagation at Frequencies Above 30 Megacycles, Proc. I.R.E., **35**, pp. 1122-1136; Oct., 1947.
3. K. Bullington, Reflection Coefficients of Irregular Terrain, Proc. I.R.E., **42**, pp. 1258-1262; Aug., 1954.
4. W. M. Sharpless, Measurements of the Angle of Arrival of Microwaves, Proc. I.R.E., **34**, pp. 837-845, Nov., 1946.
5. A. B. Crawford and W. M. Sharpless, Further Observations of the Angle of Arrival of Microwaves, Proc. I.R.E., **34**, pp. 845-848, Nov., 1946.
6. A. B. Crawford and W. C. Jakes, Selective Fading of Microwaves, B.S.T.J., **31**, pp. 68-90; January, 1952.
7. R. L. Kaylor, A Statistical Study of Selective Fading of Super High Frequency Radio Signals, B.S.T.J., **32**, pp. 1187-1202, Sept., 1953.
8. H. E. Bussey, Microwave Attenuation Statistics Estimated from Rainfall and Water Vapor Statistics, Proc. I.R.E., **38**, pp. 781-785, July, 1950.
9. J. C. Schelleng, C. R. Burrows, E. B. Ferrell, Ultra-Short Wave Propagation, B.S.T.J., **12**, pp. 125-161, April, 1933.
10. MIT Radiation Laboratory Series, L. N. Ridenour, Editor-in-Chief, Volume 13, Propagation of Short Radio Waves, D. E. Kerr, Editor, 1951, McGraw-Hill.

11. Summary Technical Report of the Committee on Propagation, National Defense Research Committee. Volume 1, Historical and Technical survey. Volume 2, Wave Propagation Experiments. Volume 3, Propagation of Radio Waves. Stephen S. Attwood, editor, Washington, D.C., 1946.
12. C. W. Burrows and M. C. Gray, The Effect of the Earth's Curvature on Ground Wave Propagation, *Proc. I.R.E.*, **29**, pp. 16-24, Jan., 1941.
13. K. Bullington Characteristics of Beyond-Horizon Radio Transmission, *Proc. I.R.E.*, **43**, p. 1175; Oct., 1955.
14. W. E. Gordon, Radio Scattering in The Troposphere, *Proc. I.R.E.*, **43**, p. 23, Jan., 1955.
15. K. Bullington, Radio Transmission Beyond the Horizon in the 40- to 4,000 MC Band, *Proc. I.R.E.*, **41**, pp. 132-135, Jan., 1953.
16. K. A. Norton, P. L. Rice and L. E. Vogler, The Use of Angular Distance in Estimating Transmission Loss and Fading Range for Propagation Through a Turbulent Atmosphere Over Irregular Terrain, *Proc. I.R.E.*, **43**, pp. 1488-1526, Oct., 1955.
17. F. H. Dickson, J. J. Egli, J. W. Herbstreit, and G. S. Wickizer, Large Reductions of VHF Transmission Loss and Fading by Presence of Mountain Obstacle in Beyond Line-Of-Sight Paths, *Proc. I.R.E.*, **41**, pp. 967-9, Aug., 1953.
18. K. Bullington, W. J. Inkster and A. L. Durkee, Results of Propagation Tests at 505 MC and 4090 MC on Beyond-Horizon Paths, *Proc. I.R.E.*, **43**, pp. 1306-1316, Oct., 1955.
19. Same as 13.
20. H. G. Booker and J. T. deBettencourt, Theory of Radio Transmission by Tropospheric Scattering Using Very Narrow Beams, *Proc. I.R.E.*, **43**, pp. 281-290, March, 1955.
21. W. H. Tidd, Demonstration of Bandwidth Capabilities of Beyond-Horizon Tropospheric Radio Propagation, *Proc. I.R.E.*, **43**, pp. 1297-1299, October, 1955.
22. K. Bullington, Radio Propagation Variations at VHF and UHF, *Proc. I.R.E.*, **38**, pp. 27-32, Jan., 1950.
23. W. R. Young, Comparison of Mobile Radio Transmission at 150, 450, 900 and 3700 MC, *B.S.T.J.*, **31**, pp. 1068-1085, Nov., 1952.
24. K. A. Norton, The Physical Reality of Space and Surface Waves in the Radiation Field of Radio Antennas, *Proc. I.R.E.*, **25**, pp. 1192-1202, Sept., 1937.
25. Same as 2.
26. C. R. Burrows, Radio Propagation Over Plane Earth-Field Strength Curves, *B.S.T.J.*, **16**, pp. 45-75, Jan., 1937.
27. K. A. Norton, The Propagation of Radio Waves Over the Surface of the Earth and in the Upper Atmosphere, Part II, *Proc. I.R.E.*, **25**, pp. 1203-1236, Sept., 1937.
28. Same as 26.
29. National Bureau of Standards Circular 462, Ionospheric Radio Propagation, Superintendent of Documents, U.S. Govt. Printing Office, Washington 25, D.C.
30. National Bureau of Standards, CRPL Series D, Basic Radio Propagation Predictions, issued monthly by U.S. Govt. Printing Office.
31. National Bureau of Standards Circular 465, Instructions for Use of Basic Radio Propagation Predictions, Superintendent of Documents, U.S. Govt. Printing Office, Washington, D.C.
32. E. W. Allen, Very-High Frequency and Ultra-High Frequency Signal Ranges as Limited by Noise and Co-channel Interference, *Proc. I.R.E.*, **35**, pp. 128-136, Feb., 1947.
33. D. K. Bailey, R. Bateman and R. C. Kirby, Radio Transmission at VHF by Scattering and Other Processes in the Lower Ionosphere, *Proc. I.R.E.*, **43**, pp. 1181-1230, Oct., 1955.
34. J. W. Herbstreit, *Advances in Electronics*, **1**, Academic Press, Inc., pp. 347-380, 1948.

A Reflection Theory for Propagation Beyond the Horizon*

By H. T. FRIIS, A. B. CRAWFORD and D. C. HOGG

(Manuscript received January 9, 1957)

Propagation of short radio waves beyond the horizon is discussed in terms of reflection from layers in the atmosphere formed by relatively sharp gradients of refractive index. The atmosphere is assumed to contain many such layers of limited dimensions with random position and orientation. On this basis, the dependence of the propagation on path length, antenna size and wavelength is obtained.

INTRODUCTION

It was pointed out several years ago¹ that power propagated beyond the radio horizon at very short wavelengths greatly exceeds the power calculated for diffraction around the earth. This beyond-the-horizon propagation has stimulated numerous experimental and theoretical investigations.² Booker and Gordon,³ Villars and Weisskopf⁴ and others have developed theories based on scattering of the radio waves by turbulent regions in the troposphere. This paper proposes a theory in which uncorrelated reflections from layers in the troposphere are assumed responsible for the power propagated beyond the horizon.

In developing this theory, some arbitrary assumptions of necessity have been made concerning the reflecting layers since, at the present time, our detailed knowledge of the atmosphere is insufficient. However, calculations based on the theory are found to be in good agreement with reported measurements of beyond-the-horizon propagation.

Measurement of the dielectric constant of the atmosphere⁵ has shown that relatively sharp variations in the gradients of refractive index exist in both the horizontal and vertical planes. Although the geometrical structure of the boundaries formed by the gradients is not well known, one may postulate an atmosphere of many layers of limited extent and

* This material was presented at the I.R.E. Canadian Convention, Toronto, Canada, October 3, 1956.

arbitrary aspect.* The number and size of the reflecting layers, as well as the magnitude of the discontinuities in the gradient of dielectric constant which form them, influence the received power.

The reflecting properties of the layers are discussed first. Next, an expression for the received power is obtained by summing the contributions of many layers in the volume common to the idealized patterns chosen to represent the transmitting and receiving antenna beams.† This expression is then used to calculate the effect on received power of changes in such parameters as the orientation of the antennas, wavelength, distance, and antenna size.

The MKS system of units is used throughout.

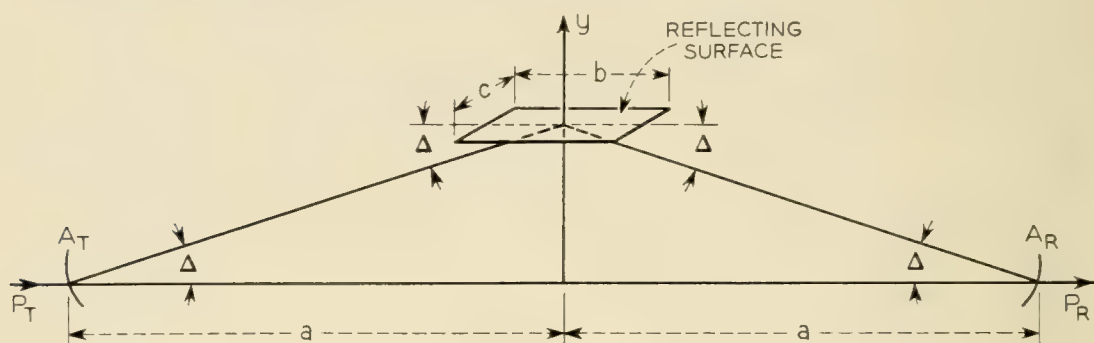


Fig. 1 — Reflection by a layer.

EFFECT OF LAYER SIZE

Propagation from a transmitting antenna of effective area A_T to a receiving antenna of effective area A_R by means of a reflecting layer is illustrated in Fig. 1. The ray from transmitter to receiver grazes the layer at angle Δ . The reflection from the layer depends on the amplitude reflection coefficient, q , which is a function of the grazing angle, and on the dimensions of the layer relative to the dimension of a Fresnel zone.‡ Three cases, depending on the layer dimensions, will be considered.

* After this paper was submitted for publication, a report was received giving some measurements of sharp variations in dielectric constant gradient and estimates of the horizontal dimensions of layers in the troposphere. J. R. Bauer, The Suggested Role of Stratified Elevated Layers in Transhorizon Short-Wave Radio Propagation, Technical Report No. 124, Lincoln Laboratory, M.I.T., Sept., 1956.

† Under some conditions, layers outside the volume common to the antenna beams may contribute appreciably to the received power. Phenomena such as multiple reflections and trapping mechanisms are not considered in this study.

‡ The power received by reflection from the layer in Fig. 1 can be calculated approximately by assuming it to be the same as the power that would be received by diffraction through an aperture in an absorbing screen, the dimensions of the aperture being the same as the dimensions of the layer projected normally to the directions of propagation. The field at the receiver is calculated from the distribution of Huygens sources in the aperture. The received power, expressed in

Case 1. Large Layers

If the layer were a plane, perfectly reflecting surface of unlimited extent, the power at the terminals of antenna A_R would be the same as the power received under line-of-sight conditions,⁶

$$P_R = P_T \frac{A_T A_R}{4\lambda^2 a^2}$$

If the layer has an amplitude reflection coefficient, q , the received power is,

$$P_R = P_T \frac{A_T A_R}{4\lambda^2 a^2} q^2$$

This relation applies when the layer dimensions are large in terms of the wavelength and are large compared with the Fresnel zone dimensions; that is, $b > \sqrt{2a\lambda}/\Delta$ and $c > \sqrt{2a\lambda}$.

Case 2. Small Layers

When the dimensions of the layer are small compared with the Fresnel zone, but large compared with the wavelength, the received power is given by the "radar" formula,

$$P_R = P_T \frac{A_T A_R}{\lambda^4 a^4} c^2 (b\Delta)^2 q^2$$

This relation applies when $b < \sqrt{2a\lambda}/\Delta$ and $c < \sqrt{2a\lambda}$.

terms of Fresnel integrals, is

$$P_R = P_T \frac{A_T A_R}{\lambda^2 a^2} [C^2(u) + S^2(u)][C^2(v) + S^2(v)]$$

where $u = \frac{c}{\sqrt{\lambda a}}$ and $v = \frac{b\Delta}{\sqrt{\lambda a}}$

When u and v are very large, we have, approximately,

$$C(u) = S(u) = C(v) = S(v) = \frac{1}{2}$$

and the expression for P_R reduces to that given for Case 1 above, except for the factor q^2 .

When both u and v are very small, we have approximately,

$$\begin{aligned} C(u) &= u & C(v) &= v \\ S(u) &= 0 & S(v) &= 0 \end{aligned}$$

and the expression for P_R reduces to that given for Case 2.

When u is large and v is small the expression for P_R given in Case 3 results.

Case 3. Layers of Intermediate Size

If the layer dimensions are such that c is large but $b\Delta$ is small, compared with the Fresnel zone dimension, the received power is given by

$$P_R = P_T \frac{A_T A_R}{2\lambda^3 a^3} (b\Delta)^2 q^2$$

In the atmosphere, c and b are likely to be about equal, on the average, and we have for this case, $\sqrt{2a\lambda} < b < \sqrt{2a\lambda}/\Delta$.

All three of these cases may be present at various times, since the structure of the atmosphere changes from day to day. However, for the purpose of the present study, Case 3 is considered most prevalent and is assumed in all the calculations to follow.

Many of the numerous layers that are assumed to contribute to the received power are not necessarily horizontally disposed, they may be oriented in any direction. Therefore, reflection in the direction of the receiver can take place from layers located both on and off the great circle path. If there are N contributing layers per unit volume in the region V common to the radiation patterns of the transmitting and receiving antennas, then for Case 3,

$$P_R = P_T \frac{A_T A_R N b^2}{2\lambda^3 a^3} \int_V \Delta^2 q^2 dV \quad (1)$$

In this relation it has been assumed that the layer size and the number of layers per unit volume remain sensibly constant throughout the common volume.

The integration process requires expressions for the reflection coefficient q and the grazing angle Δ of the layers in the common volume. These quantities are derived in the following sections.

REFLECTION COEFFICIENT OF A LAYER

The reflection coefficient of a plane boundary (Fig. 2) separating two media whose dielectric constants, relative to free space, differ by an increment $d\epsilon$ is given by Fresnel's laws of reflection. For both polarizations, the plane wave reflection coefficient of the boundary is

$$q = d\epsilon/4\Delta^2$$

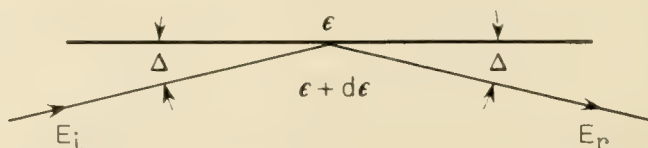


Fig. 2 — Reflection at a boundary between two homogeneous media.

provided $1 \gg \Delta^2 \gg d\epsilon$. This reflection coefficient for an incremental change in dielectric constant can be used to calculate the reflection from discontinuities in the gradient of the dielectric constant of the atmosphere such as those shown for a stratified medium at $y = 0$ and $y = h$ in Fig. 3(b).

Such variations of dielectric constant are assumed to be representative of discontinuities in gradient as they exist in the physical atmosphere. The variations form the reflecting layers.

The method of calculating the reflection coefficient of such a stratified medium is due to S. A. Schelkunoff⁷ and is illustrated schematically in Fig. 3 in which the medium has been subdivided into incremental steps. Consider the reflected wave from a typical incremental layer, dy , situated a distance y above the lower boundary of the layer, 0. From Fig. 3(a) it is clear that the phase of this wave is $4\pi y/\lambda \sin \Delta$ relative to that of a wave reflected from the lower boundary. The incremental reflection coefficient is $d\epsilon/4\Delta^2 = -K dy/4\Delta^2$, where K is the change in gradient of the dielectric constant at the boundaries of the layer. The field reflected by layer dy is therefore,

$$dE_r = -E_i \frac{K}{4\Delta^2} e^{-j(4\pi y/\lambda) \sin \Delta} dy$$

One now obtains the complete reflected field by summing the reflections from all increments within the layer of thickness h .

$$E_r = \int_0^h dE_r = jE_i \frac{K\lambda}{16\pi\Delta^2 \sin \Delta} [1 - e^{-j(4\pi h/\lambda) \sin \Delta}]$$

This relation shows that the layer is equivalent to two boundaries at

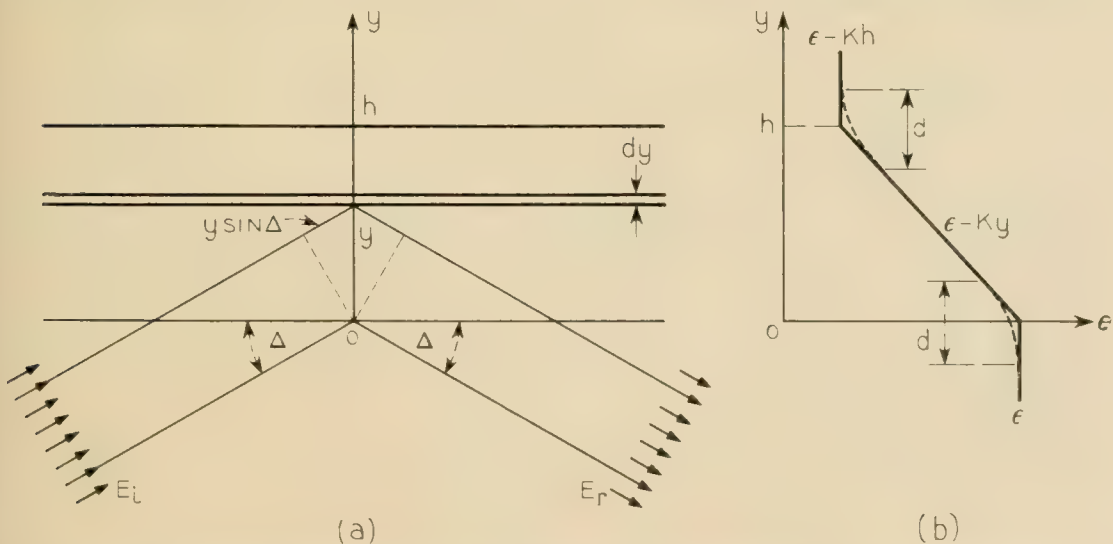


Fig. 3 — Plane-wave reflection at an incremental layer dy within a stratified medium extending from $y = 0$ to $y = h$.

$y = 0$ and $y = h$, each with reflection coefficient

$$q = \frac{K\lambda}{16\pi\Delta^3} \quad (2)$$

If the abrupt change in slope, the solid line in Fig. 3(b), is replaced by a gradual change as indicated by the dotted lines, (2) still holds provided $d < \lambda/4\Delta$. For more gradual changes, $d = n\lambda/4\Delta$ where $n > 1$, the reflection coefficient is

$$q = \frac{K\lambda}{16\pi\Delta^3} \cdot \frac{\sin \frac{\pi n}{2}}{\frac{\pi n}{2}}$$

and q varies with n between $q = 0$ and

$$q = \frac{K\lambda}{16\pi\Delta^3} \cdot \frac{2}{\pi n}$$

Smoothing of the boundaries reduces the value of q .

It will be assumed in all the calculations to follow that reflection from layers in the troposphere is described by (2).

VARIATION OF STRENGTH OF LAYERS WITH HEIGHT

The formula for the reflection coefficient includes the factor K , which represents the change in the gradient of the dielectric constant at the boundaries of the layer. A dielectric constant profile constructed of many randomly positioned gradients is shown schematically in Fig. 4. The variations are shown as departures from the standard linear gradient. Measurements⁵ indicate that the fluctuations of the dielectric constant normally decrease with height above ground. The changes in the dielectric constant gradients associated with these fluctuations probably vary in a similar manner so that K is some inverse function of the height above the earth. However, to simplify the computation of received power, to be described later, we have adopted the cylindrical coordinate system shown in Fig. 5, and it is convenient, then, to let K be a function of ρ , the distance from the chord joining the transmitter and receiver to the point in question.

We assume, therefore, that

$$K = \frac{3,200K_1}{\rho} \quad (3)$$

where K_1 is the change in gradient at point A in Fig. 5 which, for a typi-

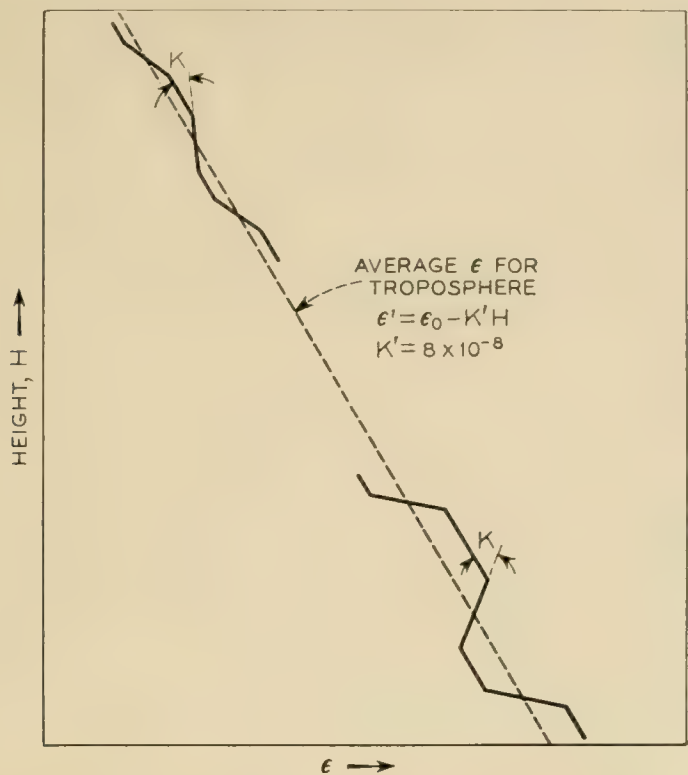


Fig. 4 — Schematic illustration of variation of the dielectric constant in the troposphere.

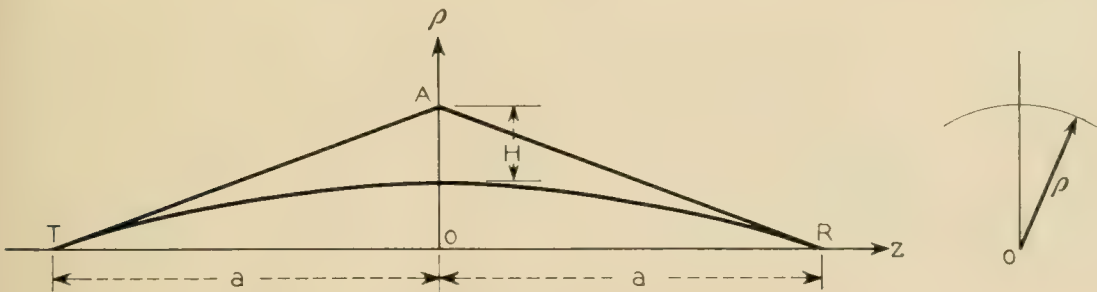


Fig. 5 — Coordinate system for a beyond-the-horizon circuit.

cal path length of 200 miles, is about 1,600 meters (1 mile) above the earth or, since $\rho = 2H$, 3,200 meters from the z axis.

Equation 3 is used in all the calculations to follow.

THE GRAZING ANGLE Δ

The grazing angle Δ at the slightly tilted layer shown in Fig. 6 is given by

$$\tan 2\Delta = \frac{2a\rho}{a^2 - z^2 - \rho^2}$$

Throughout the volume common to the antenna patterns, $\Delta \ll 1$, $\rho \ll a$ and $z < a/2$. Then

$$\Delta \approx \frac{\rho}{a} \quad (4)$$

It is evident that Δ is constant and equal to ρ/a when the point (ρ, z) is located on a cylinder with axis TR and radius ρ . It is this feature that

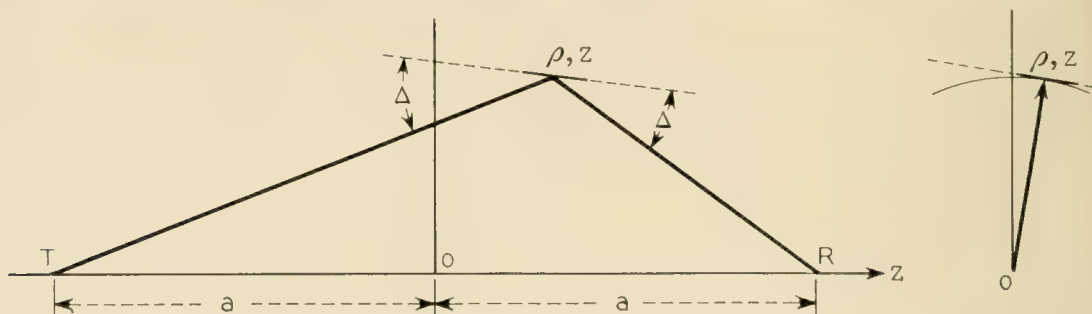


Fig. 6 — Grazing angle Δ at a layer.

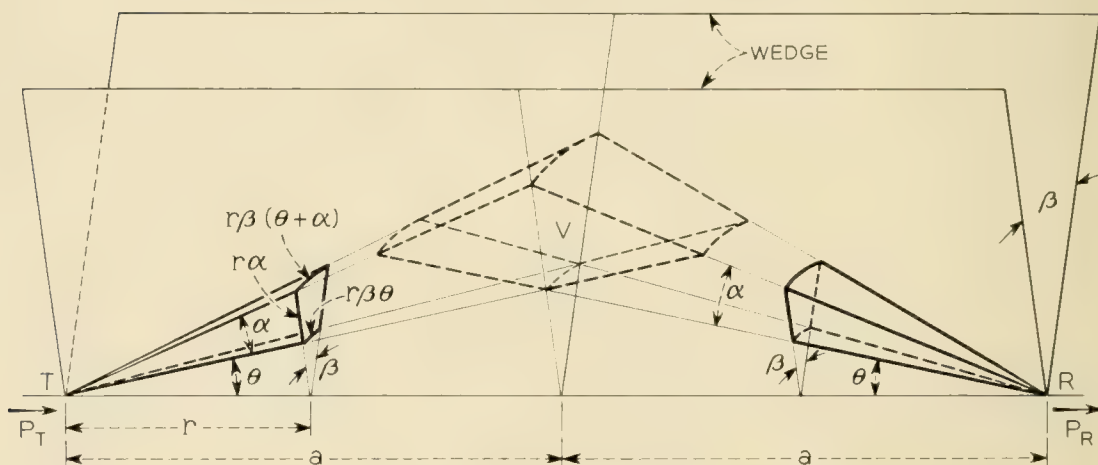


Fig. 7 — Idealized antenna patterns used in this study.

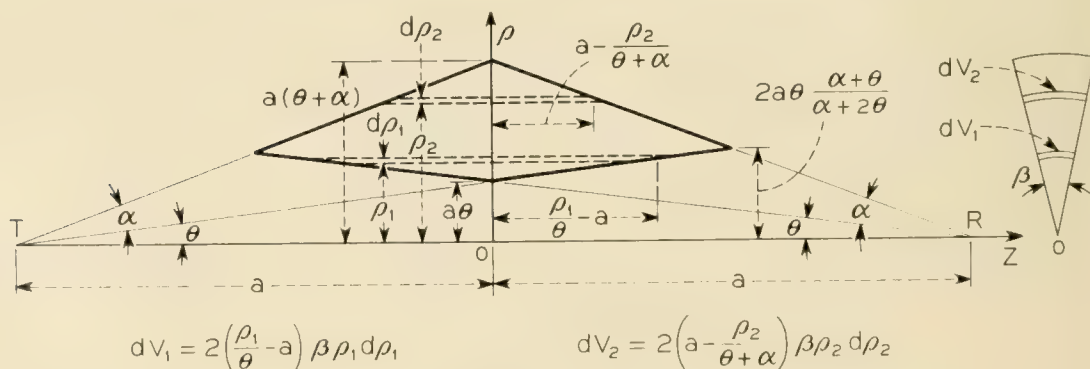


Fig. 8 — Integration over the common volume.

suggested the unusual idealized antenna patterns shown in Fig. 7, which are described in the next section.

CALCULATION OF THE RECEIVED POWER

Substituting (2), (3) and (4) in (1), one obtains for the received power,

$$P_R = P_T M a A_R A_T \lambda^{-1} \int_V \rho^{-6} dV \tag{5}$$

where

$$M \approx 2000 b^2 K_1^2 N \tag{6}$$

To integrate over the volume common to actual antenna patterns would be difficult. We have, as mentioned before, replaced the actual patterns with the idealized patterns shown in perspective in Fig. 7 and in plane projection in Fig. 8. The patterns (Fig. 7) are bounded by side planes of the large wedge and by surfaces of cones with axis TR . The common volume is indicated by broken lines and is well defined. Since the grazing angle Δ is constant for the incremental cylindrical volumes dV_1 and dV_2 shown in Fig. 8, it is easy to integrate over the common volume V and we obtain

$$\int_V \rho^{-6} dV = \frac{\beta}{6\theta^4 a^3} f\left(\frac{\alpha}{\theta}\right) \tag{7}$$

$$f\left(\frac{\alpha}{\theta}\right) = 1 + \frac{1}{\left(1 + \frac{\alpha}{\theta}\right)^4} - \frac{1}{8} \left(\frac{2 + \frac{\alpha}{\theta}}{1 + \frac{\alpha}{\theta}}\right)^4 \tag{8}$$

The function $f(\alpha/\theta)$ is plotted in Fig. 9.

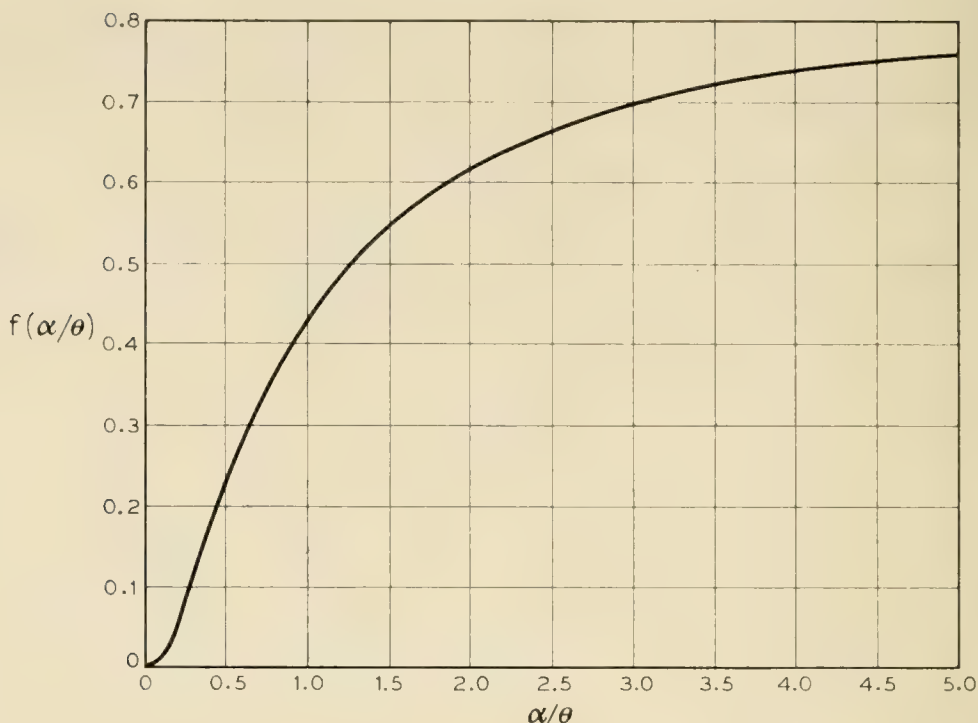
The gain of the idealized antennas is $G = 8\pi/\alpha\beta(\alpha + 2\theta)$ and the effective area is

$$A = \frac{2\lambda^2}{\alpha\beta(\alpha + 2\theta)} \tag{9}$$

The area of a cross section of the antenna pattern is bounded by two straight sides, $r\alpha$, and two curved sides $r\theta\beta$ and $r(\theta + \alpha)\beta$. The aspect ratio is defined as the ratio of the sum of the lengths of the curved sides to the sum of the lengths of the straight sides. It is equal to one when

$$\beta = \frac{2\alpha}{\alpha + 2\theta} \tag{10}$$

Substituting (10) in (9) gives the effective area of the idealized antenna

Fig. 9 — The function $f(\alpha/\theta)$.

with aspect ratio one,

$$A = \frac{\lambda^2}{\alpha^2} \quad (11)$$

Substituting (7), (10) and (11) in (5) gives for *identical transmitting and receiving antennas with aspect ratio one*,

$$P_R = P_T \frac{M \lambda^3}{3} \frac{1}{\alpha^3} \frac{1}{\theta^5 a^2} \frac{1}{2 + \frac{\alpha}{\theta}} f\left(\frac{\alpha}{\theta}\right)^* \quad (12)$$

For actual antennas, α may be taken to be the half-power beam-width.

In the following sections, (12) will be used to derive some of the general properties of propagation beyond the horizon.

* K. Bullington has suggested that a useful form for equation (12) is

$$P_R = \left[\frac{P_T \lambda^2}{4a^2 \alpha^4} \right] \left[\frac{4M\lambda}{3\theta^4} \right] \left[\frac{\frac{\alpha}{\theta} f\left(\frac{\alpha}{\theta}\right)}{2 + \frac{\alpha}{\theta}} \right]$$

where the first term in brackets represents the power that would be received in free space, the second term involves the characteristics of the troposphere, and the third term is a correction factor for narrow beam antennas.

RECEIVED POWER VERSUS ANGLE θ

The angle θ is the angle between the lower edge of the idealized antenna pattern and the straight line joining the terminals (see Fig. 7). The minimum value of θ is determined by the profile of the transmission path. If λ , α and a are constants, (12) can be used to calculate the ratio of the powers, P_{R1} and P_{R2} , received at two different angles, θ_1 and θ_2 ,

$$\frac{P_{R1}}{P_{R2}} = \left(\frac{\theta_2}{\theta_1}\right)^5 \frac{2 + \frac{\alpha}{\theta_2}}{2 + \frac{\alpha}{\theta_1}} \frac{f\left(\frac{\alpha}{\theta_1}\right)}{f\left(\frac{\alpha}{\theta_2}\right)} \quad (13)$$

Equation (13) shows the importance of having the angle θ as small as possible. For example, for $\theta_1 = \alpha$ and $\theta_2/\theta_1 = 1.25$, the power ratio is 3.4 (5 db). Thus in an actual circuit the antenna pattern should be close to the horizon plane.

RECEIVED POWER VERSUS WAVELENGTH

Consider a given path in which a and θ are specified. Equation (12) can be used to calculate the ratio of received powers, P_{R1} and P_{R2} , corresponding to two different wavelengths, λ_1 and λ_2 .

$$\frac{P_{R1}}{P_{R2}} = \left(\frac{\lambda_1}{\lambda_2}\right)^3 \left(\frac{\alpha_2}{\alpha_1}\right)^3 \frac{2 + \frac{\alpha_2}{\theta}}{2 + \frac{\alpha_1}{\theta}} \frac{f\left(\frac{\alpha_1}{\theta}\right)}{f\left(\frac{\alpha_2}{\theta}\right)} \quad (14)$$

where α_1 and α_2 are the beamwidths of the antenna patterns at wavelengths λ_1 and λ_2 respectively.

Case I. Equal antenna gains at the two wavelengths.

For this case, $\alpha_1 = \alpha_2$ and equation (14) reduces to

$$\frac{P_{R1}}{P_{R2}} = \left(\frac{\lambda_1}{\lambda_2}\right)^3 \quad (15)$$

In free space the power ratio would be

$$\frac{P_{R1}}{P_{R2}} = \left(\frac{\lambda_1}{\lambda_2}\right)^2 \quad (16)$$

or

$$\frac{P_{R1}/P_{R2} \text{ (Beyond-Horizon)}}{P_{R1}/P_{R2} \text{ (Free Space)}} = \frac{\lambda_1}{\lambda_2} \quad (17)$$

Thus if 400 mcs and 4,000 mcs were propagated over the same path, the antenna gains being identical for the two systems, then, on the average, one would expect the received power relative to the free space value at 400 mcs to be 10 db higher than that at 4,000 mcs because of the characteristics of the troposphere.

Case II. Equal antenna apertures for the two wavelengths.

For this case, $\alpha_1/\alpha_2 = \lambda_1/\lambda_2$ and (14) reduces to

$$\frac{P_{R1}}{P_{R2}} = \frac{2 + \frac{\alpha_2}{\theta}}{2 + \frac{\alpha_1}{\theta}} \frac{f\left(\frac{\alpha_1}{\theta}\right)}{f\left(\frac{\alpha_2}{\theta}\right)} \quad (18)$$

Experimental data for this case was obtained on the 150 nautical mile test circuit between St. Anthony and Gander in Newfoundland.⁸ The antennas for both wavelengths were paraboloids 8.5 meters in diameter. Simultaneous transmission tests at $\lambda_1 = 0.074$ m and $\lambda_2 = 0.6$ m were conducted for a full year. For this circuit,

$$\theta = 0.94^\circ \text{ (4/3 earth radius)}$$

$$\alpha_1 = 66 \frac{0.074}{8.5} = 0.575^\circ$$

$$\alpha_2 = 66 \frac{0.6}{8.5} = 4.65^\circ$$

Using these values in (18), we get for the ratio of received powers,

$$P_{R1}/P_{R2} = 1.01$$

For antennas of equal aperture in free space,

$$P_{R1}/P_{R2} = (\lambda_2/\lambda_1)^2 = 65.5$$

Therefore,

$$\frac{P_{R1}/P_{R2} \text{ (Beyond Horizon)}}{P_{R1}/P_{R2} \text{ (Free Space)}} = \frac{1}{65} = -18.1 \text{ db}$$

The Summary of Results, Sections 1 and 2 on page 1316 of Reference 8, gives -17 db for this ratio. The agreement between calculated and measured values is very good.

RECEIVED POWER VERSUS DISTANCE

If antenna size and wavelength are specified, (12) gives for two distances, a_1 and a_2 ,

$$\frac{P_{R1}}{P_{R2}} = \left(\frac{a_2}{a_1}\right)^7 \frac{2 + \frac{\alpha}{\theta_2}}{2 + \frac{\alpha}{\theta_1}} \frac{f\left(\frac{\alpha}{\theta_1}\right)}{f\left(\frac{\alpha}{\theta_2}\right)} \quad (19)$$

For $a_2 = 2a_1$, (19) gives for different values of α/θ_1

$\alpha/\theta_1 = 0.5$	1	2	4
$P_{R1}/P_{R2} = 276$ (24 db)	197 (23 db)	138 (21.5 db)	104 (20 db)

Fig. 1 in Bullington's paper,⁹ which gives the median signal level in decibels below the free space value as a function of distance, shows an 18 db increase in attenuation when the distance is doubled. This corresponds to a ratio of received powers of $18 + 6 = 24$ db. The examples in the table above give an average increase in attenuation of 22 db.

RECEIVED POWER VERSUS ANTENNA SIZE

Equation 14 can be used to calculate the effect on received power of changing simultaneously the size (and, hence, the beamwidths) of the antennas used for transmitting and receiving, the wavelength and distance remaining fixed.

$$\frac{P_{R1}}{P_{R2}} = \left(\frac{\alpha_2}{\alpha_1}\right)^3 \frac{2 + \frac{\alpha_2}{\theta}}{2 + \frac{\alpha_1}{\theta}} \frac{f\left(\frac{\alpha_1}{\theta}\right)}{f\left(\frac{\alpha_2}{\theta}\right)} \quad (20)$$

where P_{R1} and P_{R2} are the received powers corresponding to the antenna beamwidths α_1 and α_2 respectively.

As an example, let α_2 be constant and equal to 4° and let θ be 1° , corresponding to a 200-mile circuit. The table below gives the ratio P_{R1}/P_{R2} as α_1 is varied.

α_1	4°	2°	1°	0.5°	0.25°
P_{R1}/P_{R2} (db)	0	10	18.5	25.7	31.4
Change in db	10	8.5	7.2	5.7	

Since α is inversely proportional to the antenna dimensions, the table shows that continued doubling of the antenna dimensions results in less and less increase in output power. The increase varies from 10 to 5.7 db

in the table. This is a characteristic feature of beyond-the-horizon propagation. In free space, doubling the antenna dimensions would result in a 12-db increase in output power.

Large antennas and high power transmitters are costly, and a proper balance between their costs requires careful studies which are outside the scope of this paper. In general, it is not believed worth while from power considerations to increase the antenna size much beyond the dimensions that correspond to a pattern angle, α , equal to angle θ .

Another factor to be considered, however, is the effect of antenna size on delay distortion in beyond-the-horizon circuits. From simple path length considerations, one concludes that the delay distortion decreases when the beamwidths of the antennas are made smaller. Therefore, delay distortion requirements may dictate antenna sizes that are not justified by power considerations alone.

SEASONAL DEPENDENCE

Both the effective earth radius, R_e , and the magnitude of the discontinuities in gradient, K_1 , are related to the season of the year. During the summer when the water vapor content of the air is high, the effective radius and the discontinuities in gradient are larger than in winter. Substituting a/R_e for θ and assigning summer and winter values for R_e and K_1 , (12) may be used to calculate the ratio of the power received in summer and in winter.

$$\frac{P_R \text{ (Summer)}}{P_R \text{ (Winter)}} = \left(\frac{K_{1S}}{K_{1W}} \right)^2 \left(\frac{R_{eS}}{R_{eW}} \right)^5 \frac{2 + \frac{\alpha R_{eW}}{a}}{2 + \frac{\alpha R_{eS}}{a}} \frac{f\left(\frac{\alpha R_{eS}}{a}\right)}{f\left(\frac{\alpha R_{eW}}{a}\right)} \quad (21)$$

$$\approx \left(\frac{K_{1S}}{K_{1W}} \right)^2 \left(\frac{R_{eS}}{R_{eW}} \right)^6$$

For example, if we assume $K_{1S} = 2K_{1W}$ and $R_{eS} = 1.2 R_{eW}$, then $P_{RS}/P_{RW} = 11.9$ (10.75 db). A seasonal variation has been observed.^{8, 10}

DEPENDENCE OF RECEIVED POWER ON ANTENNA ORIENTATION

The variation of received power with orientation of the antennas at the terminals of a beyond-the-horizon circuit differs considerably from that observed under line-of-sight conditions. Consider, for example, Fig. 10 which shows the beams of the transmitting and receiving antennas elevated simultaneously. The variation of received power can be calculated from (13). As an example, consider the 188-mile circuit between

Crawford Hill, N. J., and Round Hill, Mass., for which experimental data is published.¹⁰ For this circuit $\alpha = 0.65^\circ$ (3 db points) and $\theta_1 = 1^\circ$ (4/3 earth radius). The table below gives the calculated variation of received power as angle θ_2 is varied.

$\theta_2 = 1^\circ$	1.1°	1.2°	1.4°	1.6°	1.8°	2°	2.2°
$10 \log_{10} (P_{R1}/P_{R2}) = 0$	2.3	4.5	8.5	12	15	17.9	20.5

The received power versus elevation angle, $\gamma = \theta_2 - \theta_1$, is plotted in Fig. 10. The calculated and experimental curves are in good agreement.

If the beams of the antennas are steered simultaneously in the horizontal plane, Fig. 11, the calculation of the variation of received power is

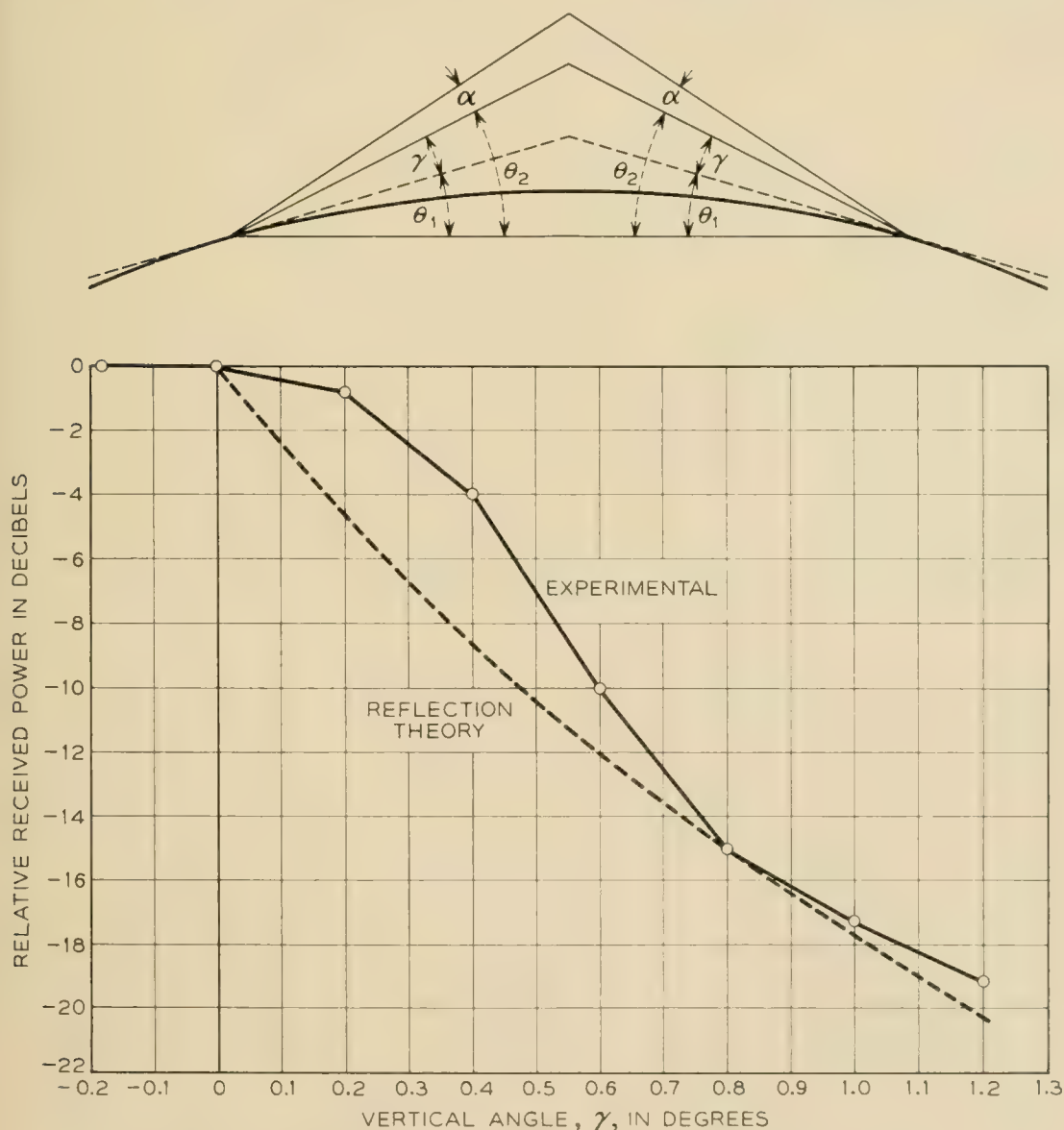


Fig. 10 — Relation between received power and vertical angle γ .

comparatively simple. In horizontal steering, the intersection of the axes of the antenna beams moves along line AB in the figure labelled "Cross section at 0." If the intersection of the beams moved along the circle $A-C$, the received power would not change. The decrease in power caused by moving the beams from position A to B is given by (13). The calcu-

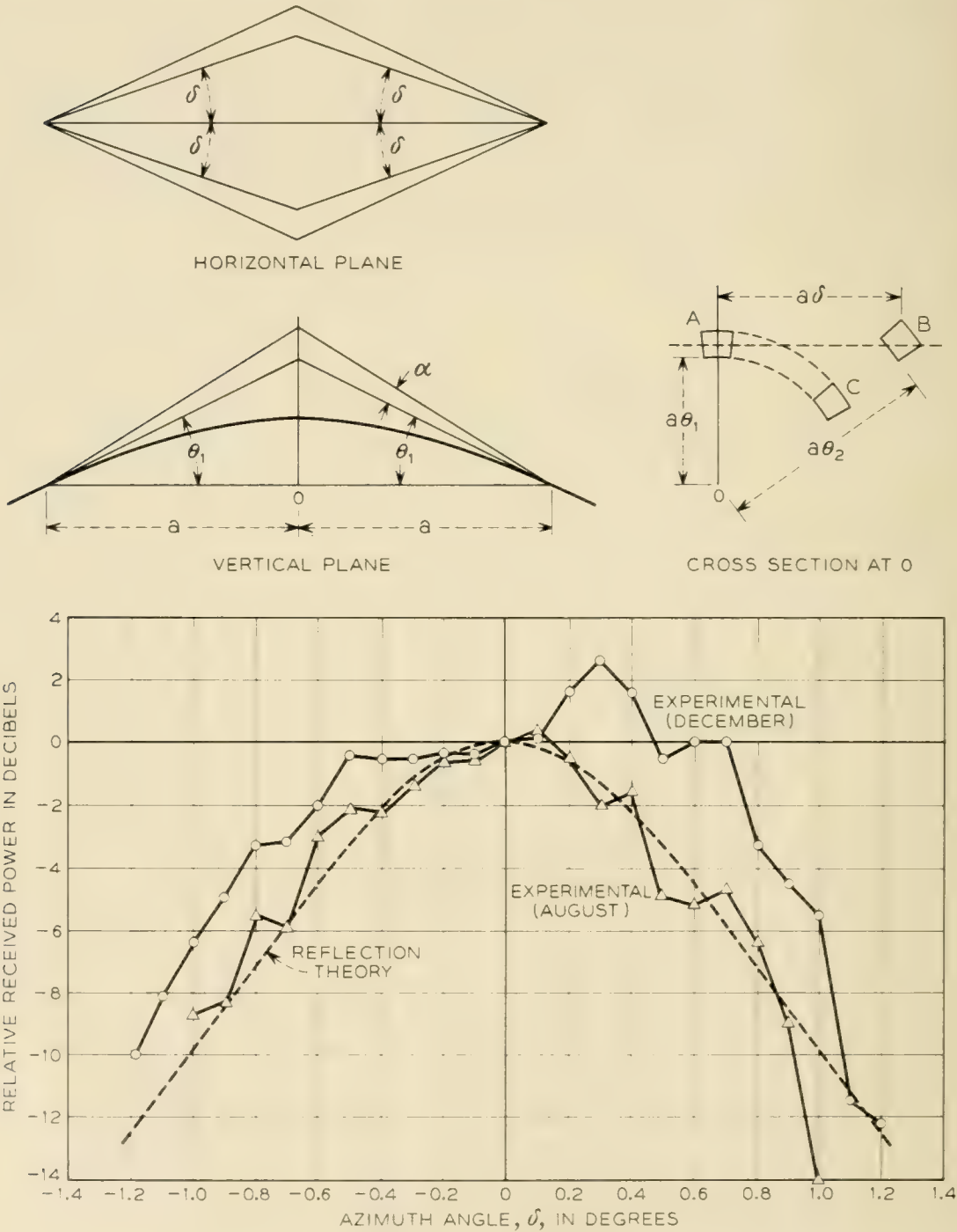


Fig. 11 — Relation between received power and azimuth angle δ .

lations are identical with the calculations for elevation steering; the elevation angle θ_2 is related to the azimuth angle δ and the beamwidth angle α by

$$\delta^2 = (\theta_2 - \theta_1)(\theta_2 + \theta_1 + \alpha)$$

A calculated curve of received power versus azimuth angle is shown in Fig. 11 for the Crawford Hill-Round Hill circuit together with the reported experimental data.¹⁰ The agreement is considered good.

THE VALUE OF FACTOR M IN EQUATION (12)

An average value for the factor M can be obtained from propagation data. Using equation (12), the ratio of received powers corresponding to free space and beyond-horizon transmission is

$$\frac{P_R \text{ (free space)}}{P_R \text{ (beyond-horizon)}} = \frac{0.75 \theta^5 \left(2 + \frac{\alpha}{\theta}\right)}{M \lambda \alpha f \left(\frac{\alpha}{\theta}\right)} \quad (22)$$

This ratio was found experimentally to be 5×10^6 (67 db) for the circuit between St. Anthony and Gander in Newfoundland.⁸ For this circuit, $\alpha = 0.081$, $\theta = 0.0164$, $\lambda = 0.6$. Substituting these values in (22) we obtain,

$$M = 3 \times 10^{-14} \quad (23)$$

Substituting this value for M in (12) leads to the following equation for a beyond-horizon tropospheric circuit,

$$P_R = P_T \times 10^{-14} \frac{\lambda^3}{\alpha^3} \frac{1}{\theta^5 a^2} \frac{1}{2 + \frac{\alpha}{\theta}} f \left(\frac{\alpha}{\theta}\right) \quad (24)$$

Equations (6) and (23) give

$$b^2 K_1^2 N = 1.5 \times 10^{-17} \quad (25)$$

Although values of the layer dimension, b , the change in gradient, K_1 , and the number of layers per unit volume, N , are not known, it is interesting to calculate N from (25) assuming reasonable values for b and K_1 .

Assuming $K_1 = 4 \times 10^{-8}$, which is half the value of K' , the average gradient of the dielectric constant in the troposphere, and $b = 1,000$ (1 km) we find $N = 10^{-8}$ or 10 layers per cubic kilometer.

CONCLUDING REMARKS

The interpretation of propagation beyond the horizon in terms of reflection from layers of limited size formed by variations in the gradient of the dielectric constant of the atmosphere leads to relatively simple results which are in good agreement with reported experimental data. The received power depends on the wavelength, the distance, and the size of the antennas used for the circuit and on the strength and size of the reflecting layers.

As mentioned earlier, the structure of the atmosphere may change markedly from time to time so that large, small and intermediate size layers play their parts at different times. Furthermore, the effective size of a given layer may be different for widely separated wavelengths, depending on the roughness of the layer in terms of the wavelength. All that can be expected of a study such as the present one is that it serve as a guide for estimating the roles of the various parameters involved in beyond-the-horizon propagation.

REFERENCES

1. K. Bullington, Radio Propagation Variations at VHF and UHF, *Proc. I.R.E.*, **38**, p. 27, Jan., 1950.
2. *Proc. I.R.E.*, Oct., 1955.
3. H. G. Booker and W. E. Gordon, A Theory of Radio Scattering in the Troposphere, *Proc. I.R.E.*, **38**, p. 401, April, 1950.
4. F. Villars and V. F. Weisskopf, Scattering of EM Waves by Turbulent Atmospheric Fluctuations, *Phys. Rev.*, **94**, p. 232, April, 1954.
5. C. M. Crain, Survey of Airborne Refractometer Measurements, *Proc. I.R.E.*, **43**, p. 1405, Oct., 1955. H. E. Bussey and G. Birnbaum, Measurement of Variation in Atmospheric Refractive Index with an Airborne Microwave Refractometer, *N.B.S. Jour. Res.*, **51**, pp. 171-178, Oct., 1953.
6. H. T. Friis, A Note on a Simple Transmission Formula, *Proc. I.R.E.*, **34**, pp. 254-56, May, 1946.
7. S. A. Schelkunoff, *Applied Mathematics for Engineers and Scientists*, D. Van Nostrand Co., Inc., p. 212, 1948, and Remarks Concerning Wave Propagation in Stratified Media, *Communication on Pure and Applied Mathematics*, **4**, pp. 117-128, June, 1951. See also, H. Bremmer, The W.K.B. Approximation as the First Term of a Geometric-Optical Series, *Communication on Pure and Applied Mathematics*, **4**, pp. 105-115, June, 1951.
8. K. Bullington, W. J. Inkster and A. L. Durkee, Results of Propagation Tests at 505 mc and 4,090 mc on Beyond-Horizon Paths, *Proc. I.R.E.*, **43**, pp. 1306-1316, Oct., 1955.
9. K. Bullington, Characteristics of Beyond-the-Horizon Radio Transmission, *Proc. I.R.E.*, **43**, pp. 1175-1180, Oct., 1955.
10. J. H. Chisholm, P. A. Portmann, J. T. deBettencourt and J. F. Roche, Investigations of Angular Scattering and Multipath Properties of Tropospheric Propagation of Short Radio Waves Beyond the Horizon, *Proc. I.R.E.*, **43**, pp. 1317-1335, Oct., 1955.

Interchannel Interference Due to Klystron Pulling

By H. E. CURTIS and S. O. RICE

(Manuscript received August 26, 1956)

A source of interchannel interference in certain multichannel FM systems is the so-called "frequency pulling effect." This effect, which occurs in systems using a klystron oscillator, is produced by an impedance mismatch between the antenna and the transmission line feeding it. In this paper expressions are developed for the magnitude of the interference when the speech load is simulated by random noise.

INTRODUCTION

In a recent paper¹ the problem of interchannel interference produced by echoes in an FM system was treated. The mathematical development in that paper can be used to calculate the distortion that arises when a Klystron oscillator is connected to an antenna through a transmission line of appreciable length.

In the system we study, the composite signal wave (the "baseband signal") from a group of carrier telephone channels in frequency division multiplex is applied to the repeller of a Klystron and thereby modulates the frequency of the Klystron output wave. If the antenna does not match the transmission line perfectly, the output frequency is altered slightly by an amount proportional to the mismatch.

This effect, known as "pulling," results in intermodulation between the individual telephone channels. In this study, the composite signal will be simulated by a random noise signal of appropriate bandwidth and power. It is assumed that some particular message channel is idle; i.e., there is no noise energy in the corresponding frequency band (which is relatively narrow in comparison with the bandwidth of the composite signal). If the system were perfect, no power would be received in this idle channel at the output of the FM detector. In the following work, the intermodulation noise falling into this channel because of the "pulling effect" will be computed. This leads to "Lewin's integral," so called, which is tabulated herein.

PULLING EFFECT

In a perfect FM system the carrier wave can be written

$$E_0(t) = A \sin [pt + \varphi(t)] \quad (1)$$

where A is a constant and the signal is $S(t) = d\varphi/dt = \varphi'(t)$, measured in radians/second. As mentioned in the Introduction, we assume that when the FM oscillator is connected directly to a transmission line with a slightly mismatched antenna at the far end, its frequency is changed. The reactive component of the input impedance of the line "pulls" the frequency of the oscillator to its new value. When the antenna is perfectly matched, there is no change in oscillator frequency.

If the characteristic impedance of the line is Z_K and the impedance of the antenna is Z_R , the impedance Z looking into the line is

$$\begin{aligned} Z &= Z_K \frac{Z_R + Z_K \tanh P}{Z_K + Z_R \tanh P} \\ &= Z_K \frac{1 + \rho e^{-2P}}{1 - \rho e^{-2P}} \end{aligned}$$

where ρ is the reflection coefficient

$$\rho = \frac{Z_R - Z_K}{Z_R + Z_K}$$

and P is the propagation constant of the line. If the loss of the line is negligible and the reflection coefficient is small, the input impedance is approximately

$$Z \doteq Z_K[1 + 2\rho(\cos \omega T - i \sin \omega T)] \text{ ohms}$$

where ω is the oscillator frequency in radians per second and T is twice the delay of the line.

It will be observed that the magnitude of the reactive component of Z oscillates as the phase angle ωT increases.

The dependence of the frequency of an oscillator upon the load reactance has been expressed by earlier workers as a "pulling figure." This figure is customarily defined as the difference between the maximum and minimum frequencies observed when the load reactance is varied over one cycle of its oscillation (the variation being accomplished, say, by increasing T). The load is taken to be such that it causes a voltage standing wave ratio of 1.5. This corresponds to a reflection coefficient of 0.20 and 14 db return loss.

In our work, we assume that the change in frequency is directly pro-

portional to the reactive component of the input impedance. More precisely, we assume that the ideal transmitter frequency of $p + \varphi'(t)$ radians/sec is changed by the pulling effect to

$$p + \varphi'(t) + 2\pi r \sin [T(p + \varphi'(t))] \text{ radians/sec} \quad (2)$$

where r is given by

$$r = 2.5 |\rho| \times (\text{Pulling Figure in cycles/sec})$$

POWER SPECTRUM OF INTERCHANNEL INTERFERENCE

The distortion produced by the pulling effect is given by the third term in (2). This distortion will be denoted by $\theta'(t)$:

$$\theta'(t) = 2\pi r \sin [pT + T\varphi'(t)] \quad (3)$$

Our problem is to compute the power spectrum of $\theta'(t)$. In particular, we are interested in the case where the signal $\varphi'(t)$ represents the composite signal wave from a group of carrier telephone channels in frequency division multiplex. All of the channels except one are assumed to be busy. Although the power spectrum of $\varphi'(t)$ is zero for frequencies in the idle channel, the same is not true for the power spectrum of $\theta'(t)$. In fact, the interchannel interference (as observed in the idle channel) is given by that portion of the power spectrum of $\theta'(t)$ which lies within the idle channel. We shall denote the corresponding interchannel interference power in the idle channel by $w_c(f) df$ where the idle channel is assumed to be of infinitesimal width and to extend from frequency $f - df/2$ to $f + df/2$. The function $w_c(f)$ will now be computed by using the procedure developed in Reference 1.

The first step is to assume the signal $\varphi'(t)$ to be a random noise current. In order to avoid writing φ' a great many times we shall set $\varphi'(t) = S(t)$, where now $S(t)$ stands for the signal. Then the autocorrelation function for the distortion $\theta'(t)$ is

$$\begin{aligned} R_{\theta'}(\tau) &= \text{avg} [\theta'(t)\theta'(t + \tau)] \\ &= (2\pi r)^2 \text{avg} [\sin (Tp + T\varphi'(t)) \sin (Tp + T\varphi'(t + \tau))] \\ &= (2\pi r)^2 \text{avg} [\sin (Tp + TS(t)) \sin (Tp + TS(t + \tau))] \\ &= \frac{(2\pi r)^2}{2} \text{avg} [\cos (TS(t) - TS(t + \tau)) \\ &\quad - \cos (2pT + TS(t) + TS(t + \tau))] \\ &= \frac{1}{2}(2\pi r)^2 \{ \exp [-T^2 R_s(0) + T^2 R_s(\tau)] \\ &\quad - \cos (2pT) \exp [-T^2 R_s(0) - T^2 R_s(\tau)] \} \end{aligned} \quad (4)$$

where

$$R_s(\tau) = \int_0^\infty w_s(f) \cos 2\pi f\tau df \quad (5)$$

and $w_s(f)$ is the power spectrum of the applied signal $S(t)$. The last expression in (4) follows from the next to the last by analogy with equation (1.14) of Reference 1.

The dc component of the distortion $\theta'(t)$ is its average value $\bar{\theta}'$ which may be computed from

$$\bar{\theta}'^2 = R_{\theta'}(\infty) = \frac{(2\pi r)^2}{2} [e^{-T^2} R_s^{(0)} (1 - \cos 2pT)] \quad (6)$$

This follows from (4) since $R_s(\infty) = 0$.

The auto-correlation function of the distortion, excluding the dc component, is then

$$R_{\theta', -\bar{\theta}'} = \frac{(2\pi r)^2}{2} [e^{-T^2} R_s^{(0)}] [(e^{T^2} R_s(\tau) - 1) - (e^{-T^2 R_s(\tau)} - 1) \cos 2pT] \quad (7)$$

The interchannel interference spectrum is

$$w_c(f) = 4 \int_0^\infty R_c(\tau) \cos 2\pi f\tau d\tau \quad (8)$$

where, by analogy with equation (1.22) of Reference 1,

$$R_c(\tau) = \frac{(2\pi r)^2}{2} [e^{-T^2} R_s^{(0)}] [(e^{T^2 R_s(\tau)} - T^2 R_s(\tau) - 1) - (e^{-T^2 R_s(\tau)} + T^2 R_s(\tau) - 1) \cos 2pT] \quad (9)$$

As mentioned before, the function $w_c(f)$ is of interest because

$$P_I = w_c(f) df \quad (10)$$

is the average interference power appearing at the receiver in an idle channel of width df centered on frequency f .

RATIO OF INTERCHANNEL INTERFERENCE TO SIGNAL POWER

The average signal power appearing in a busy channel of width df centered on the frequency f is

$$P_s = w_s(f) df \quad (11)$$

and hence the ratio of the interchannel interference power to the signal

power is

$$\frac{P_I}{P_s} = \frac{w_c(f)}{w_s(f)} \quad (12)$$

We now obtain an expression for this ratio on the assumption that the random noise signal $S(t)$ (which is used to simulate the multichannel signal) has the power spectrum

$$w_s(f) = \begin{cases} P_0, & 0 < f < f_b \\ 0, & f > f_b \end{cases} \quad (13)$$

where P_0 is a constant. $S(t)$ is measured in radians/sec and $P_0 f_b$ is measured in (radians/sec)². $P_0 f_b$ is given by

$$P_0 f_b = \overline{S^2(t)} = \text{avg} [\varphi'(t)]^2 = (2\pi\sigma)^2$$

where σ is the rms frequency deviation of the signal measured in cycles/second. According to (5) this signal has the auto-correlation function

$$R_s(\tau) = \int_0^{f_b} P_0 \cos 2\pi f \tau \, df = P_0 \left[\frac{\sin 2\pi f \tau}{2\pi \tau} \right]_0^{f_b} = (2\pi\sigma)^2 \frac{\sin u}{u} \quad (14)$$

where

$$u = 2\pi f_b \tau$$

The interference power spectrum $w_c(f)$ corresponding to the $w_s(f)$ of (13) may be obtained by substituting (14) in (9) to get $R_c(\tau)$ and then using (8). The result is

$$\begin{aligned} w_c(f) = 4 \frac{(2\pi r)^2}{2} e^{-b} \int_0^\infty [(e^{bu^{-1} \sin u} - bu^{-1} \sin u - 1) \\ - (e^{-bu^{-1} \sin u} + bu^{-1} \sin u - 1) \cos 2pT] \frac{\cos au}{2\pi f_b} du \end{aligned} \quad (15)$$

where u is the same as in (14) and we have set

$$a = f/f_b \quad b = (2\pi\sigma T)^2$$

This integral may be expressed in terms of Lewin's integral which is studied in Appendix III of Reference 1. Thus

$$w_c(f) = \frac{(2\pi r)^2 e^{-b}}{2\pi f_b} [I(b, a) - I(-b, a) \cos 2pT] (\text{radian/sec})^2 / \text{cps} \quad (16)$$

where $I(b, a)$ and $I(-b, a)$ are tabulated for various values of a and b . Since we began the problem by dealing directly with $\theta'(t)$ which is a radian frequency, rather than $\theta(t)$ which is a radian phase, $w_c(f)$ has the

TABLE I—VALUES OF $e^{-b}I(b, a)$ FOR $b > 0$

b	e^b	$e^{-b}I(b, a)$					
		$a = 0$	0.25	0.50	0.75	1.00	1.25
0.0	1.000	0.000	0.000	0.000	0.000	0.000	0.000
0.25	1.284	0.082	0.072	0.062	0.052	0.042	0.031
0.5	1.649	0.272	0.241	0.209	0.176	0.142	0.107
1.0	2.718	0.761	0.685	0.602	0.511	0.414	0.314
2.0	7.389	1.560	1.440	1.291	1.117	0.919	0.713
3.0	20.08	1.913	1.801	1.645	1.448	1.215	0.968
4.0	54.60	1.974	1.888	1.751	1.566	1.341	1.098
5.0	148.4	1.905	1.844	1.731	1.571	1.372	1.153
6.0	403.4	1.794	1.751	1.660	1.525	1.356	1.166
7.0	1097.	1.680	1.649	1.575	1.463	1.320	1.157
8.0	2981.	1.576	1.552	1.492	1.398	1.277	1.138

TABLE II—VALUES OF $I(b, a)$ FOR $b < 0$

b	$I(b, a)$					
	$a = 0$	0.25	0.50	0.75	1.0	1.25
0.0	0.000	0.000	0.000	0.000	0.000	0.000
-0.25	0.092	0.080	0.068	0.057	0.045	0.034
-0.5	0.349	0.300	0.254	0.210	0.167	0.125
-1.0	1.25	1.06	0.885	0.723	0.576	0.432
-2.0	4.16	3.41	2.76	2.20	1.76	1.34
-3.0	8.03	6.37	4.97	3.88	3.14	2.46
-4.0	12.6	9.66	7.23	5.49	4.55	3.74
-5.0	17.8	13.2	9.40	6.89	5.93	5.19
-6.0	23.6	16.8	11.4	8.00	7.25	6.85
-7.0	30.0	20.7	13.1	8.71	8.48	8.78
-8.0	37.2	24.8	14.5	8.93	9.59	11.0

dimensions of $(\text{radians/sec})^2/\text{cps}$. The signal in the same dimensions is P_0 or $(2\pi\sigma)^2/f_b$. Therefore the ratio of the interchannel interference power to the signal power is:

$$\frac{P_I}{P_s} = \frac{1}{2\pi} \left(\frac{r}{\sigma} \right)^2 e^{-b} [I(b, a) - I(-b, a) \cos 2pT] \quad (17)$$

The quantity $e^{-b}I(b, a)$ for $b > 0$ is tabulated in Table I. The quantity $I(b, a)$ for $b < 0$ is given in Table II. These tables, which are also given in Reference 1, are repeated here for the convenience of the reader.

When the rms frequency deviation σ is so small that $b = (2\pi\sigma T)^2$ is small compared to unity, the approximation

$$I(b, a) \approx b^2 \pi (2 - a) / 4$$

leads to

$$\frac{P_I}{P_s} \approx (2\pi^2 r \sigma T^2)^2 (2 - a) (1 - \cos 2pT) / 2 \quad (18)$$

When σ and T are such that $b \gg 1$, the approximation

$$I(b, a) \approx (6\pi/b)^{1/2} \exp \left[b - \frac{3a^2}{2b} \right]$$

leads to
$$\frac{P_I}{P_s} \approx \left(\frac{3}{8\pi^3} \right)^{1/2} \frac{r^2}{\sigma^3 T} \exp \left[-\frac{3}{2} \left(\frac{a}{2\pi\sigma T} \right)^2 \right]$$

Equation (17), when converted to decibels, breaks down conveniently into two terms which may be designated D_1 and D_2 :

$$10 \log P_I/P_s = D_1 + D_2$$

$$D_1 = 10 \log (r/\sigma)^2$$

$$D_2 = 10 \log \frac{1}{2\pi} e^{-b} [I(b, a) - I(-b, a) \cos 2pT] \quad (19)$$

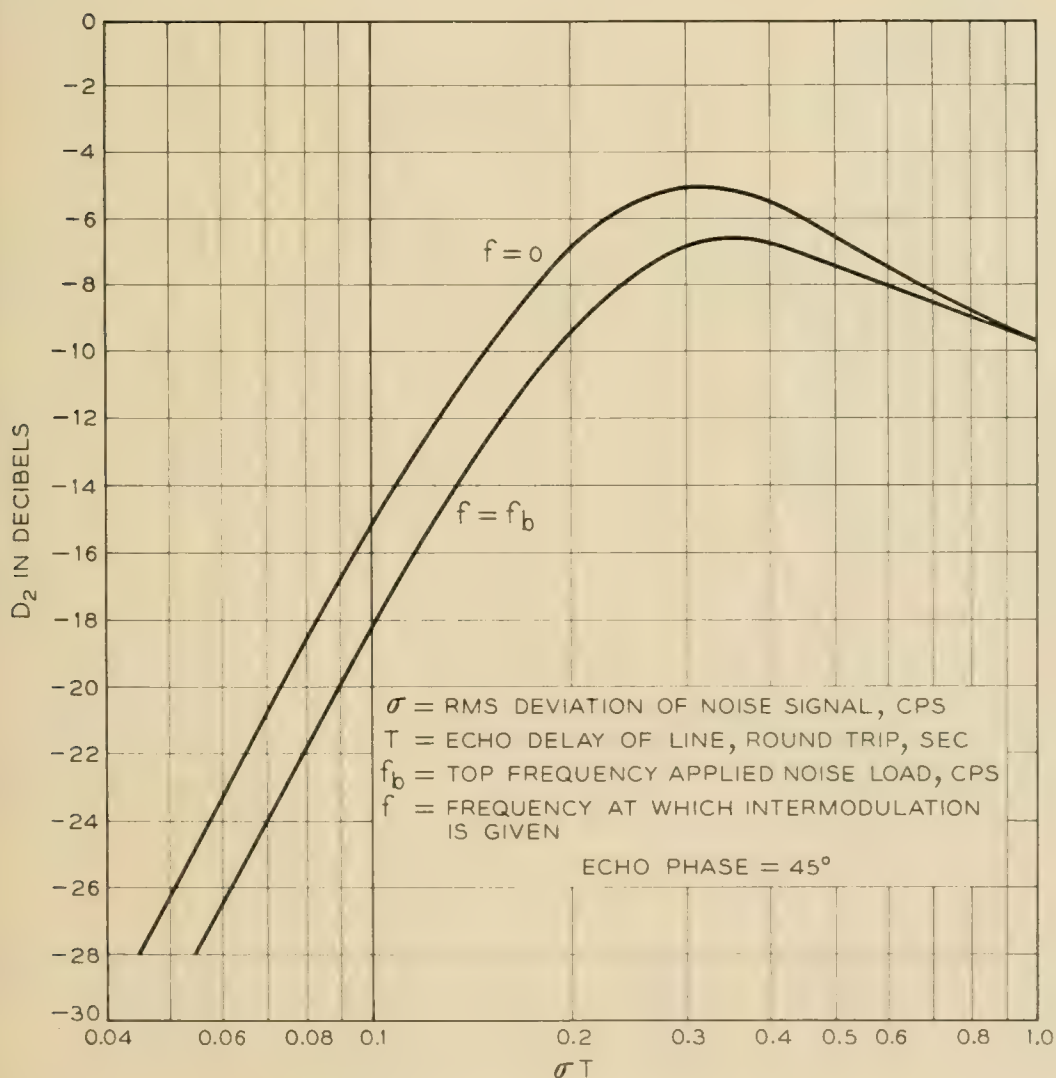


Fig. 1 — Plot of D_2 as a function of σT

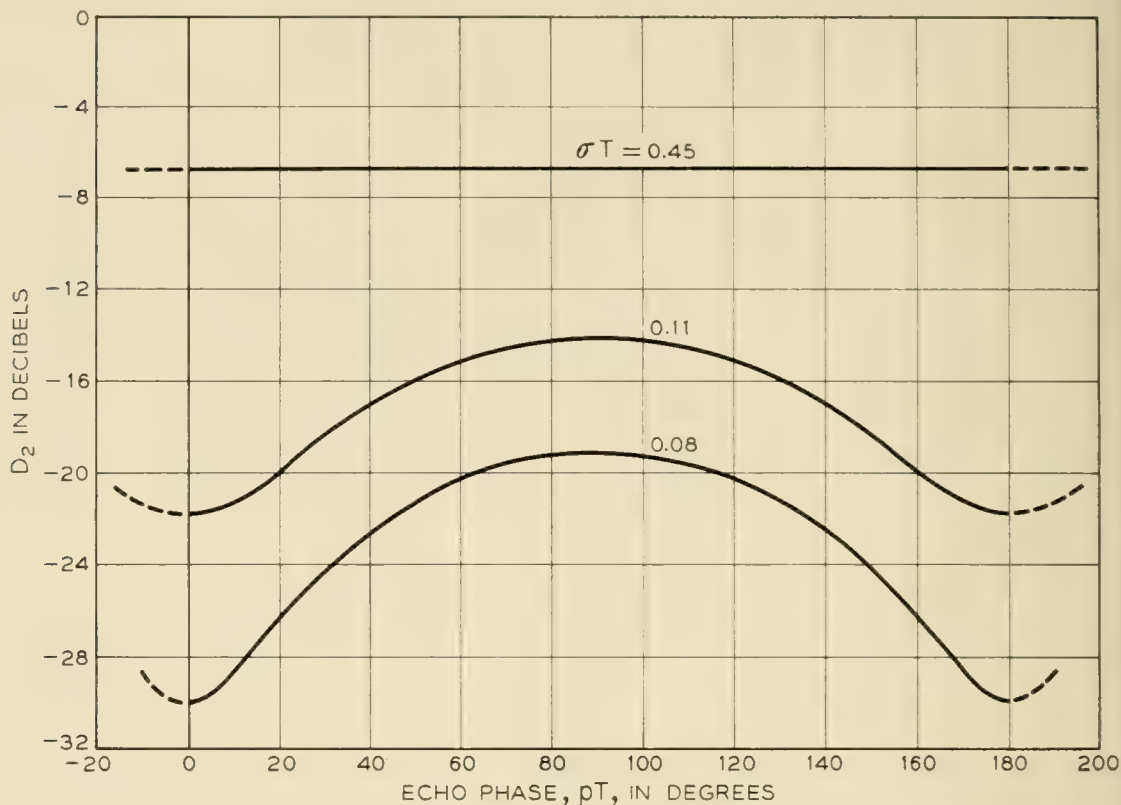


Fig. 2 — Plot of D_2 as a function of echo phase when $f = f_b$

The term D_2 depends only on the rms deviation σ , the round-trip echo delay T of the line, the ratio $a = f/f_b$, and the echo phase pT . The term D_2 is plotted in Fig. 1 as a function of σT for two channels, one at the top and the other at the bottom of the signal band. Since the carrier frequency may be expected to be very high, the carrier phase $2pT$ will be a very large number of radians even with very short wave guide runs. Hence the curves on Fig. 1 are plotted for the average value of $\cos 2pT$ which is zero.

The curves on Fig. 2 show how D_2 depends on pT and the parameter σT . In this case the channel is taken at the top of the signal band. When pT is an odd multiple of 90 degrees, it turns out that we have even order modulation products only; and when pT is a multiple of 180 degrees, odd order products only. The curves show that the distortion becomes less dependent on the echo phase as the quantity σT increases.

REFERENCE

1. W. R. Bennett, H. E. Curtis and S. O. Rice, Interchannel Interference in FM and PM Systems, B.S.T.J., **34**, pp. 601-636, May, 1955.

Instantaneous Companding of Quantized Signals

By BERNARD SMITH

(Manuscript received October 8, 1956)

Instantaneous companding may be used to improve the quantized approximation of a signal by producing effectively nonuniform quantization. A revision, extension, and reinterpretation of the analysis of Panter and Dite permits the calculation of the quantizing error power as a function of the degree of companding, the number of quantizing steps, the signal volume, the size of the "equivalent dc component" in the signal input to the compressor, and the statistical distribution of amplitudes in the signal. It appears, from Bennett's spectral analysis, that the total quantizing error power so calculated may properly be studied without attention to the detailed composition of the error spectrum, provided the signal is complex (such as speech or noise) and is sampled at the minimum information-theoretic rate.

These calculations lead to the formulation of an effective process for choosing the proper combination of the number of digits per code group and companding characteristic for quantized speech communication systems. An illustrative application is made to the planning of a hypothetical PCM system, employing a common channel compandor on a time division multiplex basis. This reveals that the calculated companding improvement, for the weakest signals to be encountered in such a system, is equivalent to the addition of about 4 to 6 digits per code group, i.e., to an increase in the number of uniform quantizing steps by a factor between $2^4 = 16$ and $2^6 = 64$.

Comparison with the results of related theoretical and experimental studies is also provided.

TABLE OF CONTENTS

(Passages marked with an asterisk contain mathematical details which may be omitted in a first reading without loss of continuity.)

	Page
I. Introduction	655
A. Fundamental Properties of Pulse Modulation	655
1. Unquantized Signals	655
2. Quantized Signals (PCM)	655
B. Quantizing Impairment in PCM Systems	656

C. Physical Implications of Nonuniform Quantization.....	657
1. Quantizing Error as a Function of Step Size.....	657
2. Properties of the Mean Square Excited Step Size.....	658
D. Nonuniform Quantization Through Uniform Quantization of a Compressed Signal.....	659
E. The Mechanism of Companding Improvement in Various Communication Systems.....	661
1. Syllabic Companding of Continuous Signals.....	661
2. Instantaneous Companding of Unquantized Pulse Signals.....	662
3. Instantaneous Companding of Quantized Signals.....	662
F. Applicability of the Present Analysis.....	663
1. Signal Spectrum.....	663
2. Sampling Rate.....	663
3. Number of Quantizing Steps.....	664
4. Subjective Effects Beyond the Scope of the Present Analysis.....	664
II. Evaluation of the Mean Square Quantization Error (σ).....	665
A.* Generalization of the Analysis of Panter and Dite.....	665
B. Operational Significance of σ	667
III. Choice of Compression Characteristic.....	667
A. Restriction to Logarithmic Compression.....	667
B. Comparison with Other Companders.....	671
IV.* The Calculation of Quantizing Error.....	672
A. Logarithmic Companding in the Absence of "DC Bias".....	672
B. Logarithmic Companding in the Presence of "DC Bias".....	673
C. Application to Speech as Represented by a Negative Exponential Distribution of Amplitudes.....	674
D. Uniform Quantization: $\mu = 0$	676
V. Discussion of General Results.....	676
A. Quantitative Description of Conventional Operation ($e_0 = 0$).....	677
1. Number of Quantizing Steps (N).....	677
2. Compandor Overload Voltage (V).....	677
3. Relative Signal Power (C).....	677
4. Average Absolute Signal Amplitude ($ e $).....	678
5. Degree of Compression (μ).....	679
B. Optimum Compandor Ensemble.....	679
C. Companding Improvement for $e_0 = 0$	682
D. Companding Improvement for $e_0 \neq 0$	684
VI. Application of Results to a Hypothetical PCM System.....	686
A. Speech Volumes.....	686
B. Choice of Compression Characteristic.....	687
1. Ideal Behavior for Speech.....	687
2. Practical Limitations on Companding Improvement.....	688
(a) Mismatch Between Zero Levels of Signal and Compandor.....	688
(b) Background Noise Level.....	690
C. Choice of the Number of Digits per Code Group.....	690
1. Ideal Behavior for Speech.....	690
2. Practical Limitations.....	693
(a) Mismatch of Zero Levels of Signal and Compandor.....	693
(b) Background Noise Level.....	695
D. Possibility of Using Automatic Volume Regulation.....	697
E. Comparison with Previous Experimental Results.....	697
VII. Conclusions.....	698
Acknowledgments.....	698
Appendix — The Minimization of Quantizing Error Power.....	698
References.....	708

I. INTRODUCTION

Quantized pulse modulation has been the subject of considerable attention in the last decade.¹⁻¹⁷ Proposals for practical application of such modulation usually provide for the transmission, by time division multiplex, of a class of signals covering an extensive power range.^{1, 6} Such proposals almost invariably assign a vital role to instantaneous companders. The present discussion is devoted to the formulation of general quantitative criteria for the choice of a suitable companding characteristic.

A. *Fundamental Properties of Pulse Modulation**

1. *Unquantized Signals*

Unquantized pulse signals are produced when a band-limited signal (such as low-pass filtered speech) is sampled instantaneously at a rate greater than or equal to the minimum acceptable value of slightly more than twice the top signal frequency. The transmission of the continuous range of pulse amplitudes so produced is known as pulse amplitude modulation (PAM). Alternatively one may translate the sampled amplitudes into variations either in the width of periodic pulses of constant amplitude (pulse duration modulation or PDM), or in the spacing of pulses of uniform amplitude and width (pulse position modulation or PPM). Regardless of the mode of transmission, the unquantized signal pulses are sensitive to noise in the transmission medium.

2. *Quantized Signals (PCM)*

Although sampling constitutes temporal quantization, it is convenient to adhere to conventional usage (as codified by Bennett² and Black¹) in restricting the designation "quantized signals" to those which have been quantized in amplitude, as well as sampled in time, in order to permit encoded (i.e., essentially telegraphic) transmission. Thus a finite range of possible signal amplitudes, large enough to accommodate the strongest signal to be encountered in a given application, may be divided into N equal parts or quantizing steps. Each instantaneous pulse amplitude of a PAM signal is then compared with this ladder-like array; amplitude quantization is accomplished by replacing all amplitudes falling in any portion of a quantizing step by a single value uniquely characterizing that interval.

* This brief account is intended merely to specify the minimum amount of background information required to avoid confusion in the present discussion. Details may be found in the many excellent and readily accessible references already cited.

Use of a binary number representation permits the encoded transmission of the N possible quantized amplitudes in terms of groups of on-off pulses containing n binary digits per code group (where $N = 2^n$).^{*} These pulses may be considered impervious to noise in the transmission medium in the sense that complete information is conveyed by the mere recognition of the presence or absence of a pulse rather than a determination of a precise magnitude. Consequently such pulses may, in principle, repeatedly be regenerated in the transmission medium, provided that regeneration occurs before the on-off pulses have been rendered indistinguishable from each other by noise.

The designation pulse code modulation (PCM) may therefore be used synonymously with quantized pulse modulation to distinguish the latter from the previously defined varieties of unquantized pulse modulation. In view of the restriction of present interest to the role of quantization *per se*, there is no need to proceed beyond the choice of quantized PAM as the prototype PCM signal in this discussion, in spite of the fact that PDM and PPM may also be quantized to yield PCM.¹

B. Quantizing Impairment in PCM Systems

From the foregoing it is clear that quantization (i.e., the representation of a bounded continuum of values by a finite number of discrete magnitudes), permits the encoded, and therefore essentially noise-free transmission of approximate, rather than exact values of sampled amplitudes. In fact, *the deliberate error imparted to the signal by quantization is the significant source of PCM signal impairment.*¹⁻⁵ Adequate limitation of this quantization error is therefore of prime importance in the application of PCM to communication systems.

A number of methods of reducing quantizing error suggest themselves on a purely qualitative and intuitive basis. For example, one may obtain a finer-grained approximation by providing more, and therefore smaller, quantizing steps for a given range of amplitudes. Alternatively, one may provide a more complete description of the signal by increasing the sampling rate beyond the minimum information-theoretic value already assumed.[†]

It is also possible to vary the size of the quantizing steps (without adding to their number) so as to provide smaller steps for weaker signals.

^{*} Of course, number representations using a base, b , other than two, so that $N = b^n$, are also available. These are presently of academic interest in view of the increased complexity of instrumentation they imply.³

[†] See Fig. 5 of Reference 2 for a quantitative evaluation of the efficacy of this measure.

Whereas the first two techniques result in an increase of bandwidth and system complexity, the third requires only a modest increase in instrumentation without any increase in bandwidth.* This investigation is therefore devoted to the study of nonuniform step size as a means of reducing quantizing impairment.

C. Physical Implications of Nonuniform Quantization

1. Quantizing Error as a Function of Step Size

Quantizing impairment may profitably be expressed in terms of the total mean square error voltage since the ratio of the mean square signal voltage to this quantity is equal to the signal-to-quantizing error power ratio.

In evaluating the mean square error voltage, we begin by considering a complex signal, such as speech at constant volume, whose pulse samples yield an amplitude distribution corresponding to the appropriate probability density. These pulse samples may be expected to fall within, or "excite", all the steps assigned to the signal's peak-to-peak voltage range. It will be assumed that, for quality telephony, the steps will be sufficiently small, and therefore numerous, to justify the assumption that the probability density is effectively constant within each step, although it may be expected to vary from step to step. Thus the continuous curve representing the probability density as a function of instantaneous amplitude is to be replaced by a suitable histogram.

If the midstep voltage is assigned to all amplitudes falling in a particular quantizing interval, the absolute value of the error in any pulse sample will be limited to values between zero and half the size of the step in question; when combined with the assumed approximation of a uniform probability density within each step, this choice minimizes the mean square error introduced at each level.⁵ Summation of the latter quantity over all levels then yields the result that the total mean square quantizing error voltage is equal to one-twelfth the weighted average of the square of the size of the voltage steps traversed (i.e., excited) by the input signal. The direct consideration of the physical meaning of this result (which, as (6) below, will constitute the basis of all subsequent calculations) will now be shown to provide a simple qualitative description of the implications of nonuniform quantization.

* We refer to bandwidth in the transmission medium as determined by the pulse repetition rate, which, in the time division multiplex applications envisioned herein, is given by the product: (sampling rate) $\times$ (number of digits or pulses per sample) $\times$ (number of channels).

2. Properties of the Mean Square Excited Step Size

Fig. 1(a) shows the range of input voltages, between the values $+V$ and $-V$ divided into N equal quantizing steps (i.e., uniformly quantized); Fig. 1(b) depicts the same range divided into N tapered steps, corresponding to nonuniform quantization.

Consider a complex signal, such as speech, whose distribution of instantaneous amplitudes at constant volume results in the excitation (symmetrically about the zero level) of the steps in the moderately large interval $X-X'$. The quantizing error power will be shown to be proportional to the (weighted) average of the square of the excited step size. For uniform quantization, it is clear, from Fig. 1(a), that this average is a constant, independent of the statistical properties of the signal. For a nonuniformly quantized signal, [Fig. 1(b)], the mean square excited step size is reduced by the division of the identical interval $X-X'$ into more steps, most of which are smaller than those shown in Fig. 1(a). Appreciation of the full extent to which the quantizing error power may so be reduced requires the added recognition that the few larger quantizing steps in the range $X-X'$, corresponding to excitation by comparatively rare speech peaks, are far less significant in their contribution to the weighted average than the small steps in the vicinity of the origin, due to the nature of the probability density of speech at constant volume.¹⁸

It is also clear that weaker signals, corresponding to a contraction of the interval $X-X'$, enjoy the greatest potential tapering advantage since their excitation may be confined to steps which are all appreciably smaller than those in Fig. 1(a). However, if the interval $X-X'$ were to

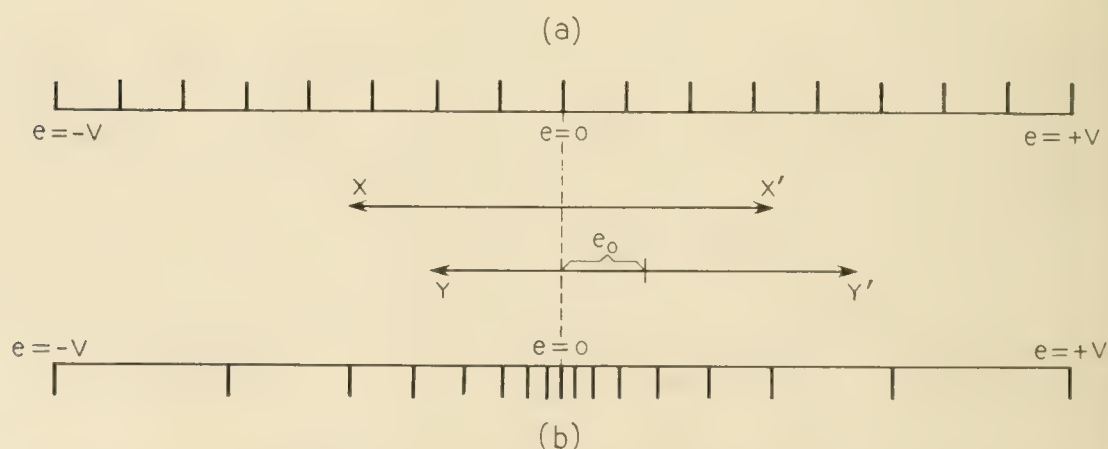


Fig. 1 — (a) Distribution of steps of equal size corresponding to direct, uniform quantization; (b) nonuniform quantization of this range into the same number of steps. The function of the instantaneous compander is to provide such nonuniform quantization in the manner illustrated in Fig. 2.

increase in size and approach the full range, $+V$ to $-V$, (to accommodate stronger signals), the excitation of extremely large steps might result in an rms step size exceeding the uniform size shown in Fig. 1(a).

Fig. 1 also indicates that signals (including unwanted noise) too weak to excite even the first quantizing step (and therefore absolutely incapable of transmission) when uniformly quantized, may successfully be transmitted as a result of the excitation of a few steps following non-uniform quantization.

Although the assumption that the average value of the signal is zero is quite proper for speech, subsequent discussion will disclose the possibility that the quiescent value of the signal, as it appears at the input to the quantizing equipment, may not always coincide with the exact center of the voltage range depicted in Fig. 1. This effect may formally be described in terms of the addition of an equivalent dc bias to the speech input at the quantizer. As shown in Fig. 1, the addition of such a dc component, e_0 , to the signal which previously excited the band of steps labeled $X-X'$, transforms $X-X'$ into an array of equal extent $Y-Y'$, centered about $e = e_0$ instead of $e = 0$. This causes the excitation of some larger steps, in Fig. 1(b), as well as the assignment of greatest weight¹⁸ to the steps in the vicinity of $e = e_0$, which are larger than those near $e = 0$; the net result is an increase in the rms excited step size, and the quantizing error power. This effect will depend on the comparative size of e_0 and the signal as well as on the degree of step size variation. In particular, Fig. 1(a) indicates that the presence of e_0 does not affect the rms excited step size under conditions of uniform quantization.

It is clear from the foregoing that the effect of nonuniform quantization of PCM signals will vary greatly with the strength of the signal; greatest improvement is to be expected for weak signals, whereas an actual impairment may be experienced by strong signals. The range of signal volumes is therefore of prime importance in the choice of the proper distribution of step sizes.

D. Nonuniform Quantization Through Uniform Quantization of a Compressed Signal

Nonuniform quantization is logically equivalent to uniform quantization of a "compressed version" of the original input signal. When applied directly, tapered quantization provides an acceptably high ratio of sample amplitude to sample error for weak pulses, by decreasing the errors (i.e., the step sizes) assigned to small amplitudes. Signal compression achieves the same goal by increasing weak pulse amplitudes without altering the step size.

The instantaneous compressor envisioned herein is, in essence, a non-linear pulse amplifier which modifies the distribution of pulse amplitudes in the input PAM signal by preferential amplification of weak samples. A satisfactory compression characteristic will have the general shape shown in Fig. 2. Thus the amplification factor, (v/e) , varies from a large value for small inputs to unity for the largest amplitude (V) to be accommodated, so that the distribution of pulse sizes may be modified without changing the total voltage range. Fig. 2 also illustrates how uniform quantization of the compressor output produces a tapered array of input steps similar to those already considered in connection with Fig. 1(b).

A complementary device, the expander, employs a characteristic inverse to that of the compressor to restore the proper (quantized) distribution of pulse amplitudes after transmission and decoding. Taken together, the compressor and expander constitute a compandor.

The resolution of tapered quantization into the sequential application of compression, uniform quantization, and expansion is operationally convenient,⁶ as well as logically sound. Since there is a one-to-one correspondence between step size allocations and compression characteristic

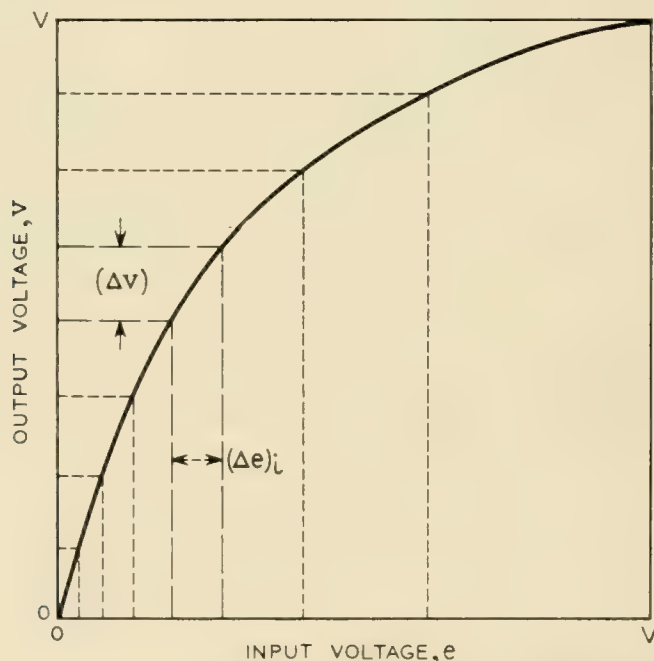


Fig. 2 — Curve illustrating the general shape of a suitable instantaneous compression characteristic. All continuous, single-valued curves connecting the origin to the point (V, V) and rising from the origin with a slope greater than one, i.e., $(dv/de)_{e=0} > 1$, are potential compression characteristics. The symmetrical negative portion $[v(e) = -v(-e)]$ is not shown. The production of a tapered array of input steps $(\Delta e)_i$ by uniform quantization of the output into steps of (equal) size Δv , is also represented.

curves, the central problem of choosing the proper distribution of step sizes will be discussed in terms of the choice of the appropriate compression characteristic; the reduction of quantizing error, corresponding to nonuniform quantization without change in the total number of steps, will be termed companding improvement.

E. *The Mechanism of Companding Improvement in Various Communication Systems*

1. *Syllabic Companding of Continuous Signals*^{19, 20}

Originally, the compandor consisted of a compressor and complementary expander operating at a syllabic rather than instantaneous rate in frequency division systems, since instantaneous companding was found to imply an undesirable increase in bandwidth in such systems.¹⁹

In spite of the existence of syllabic power variations, a useful understanding of such compandor action may be inferred from the consideration of the long-time average power. Thus, in its simplest form, the compressor might provide amplification varying from a constant value within the range of volumes corresponding to weak speech to little or no amplification for comparatively strong signals prior to transmission. Although it is an amplifying device, the compressor takes its name from the contraction of the transmitted volume range which results from selective amplification of the weakest signals. Since the distortion of the signal by the compressor may virtually be confined to a change in loudness, the compressor output may be expected to be intelligible.

In interpreting a compression characteristic, syllabic application permits the identification of the ordinate and abscissa with $\sqrt{v^2}$ and $\sqrt{e^2}$, rather than v and e as shown in Fig. 2. This substitution of rms for instantaneous signals not only confines the significance of the compression characteristic to the first quadrant but also removes the need for compandor response to input signals below some small, nonzero, threshold value.

If we designate the mean square noise voltage in the transmission medium by $\overline{v_n^2}$, the amplification of weak signals prior to exposure to this noise provides an increase in the transmitted signal to noise ratio from $(\overline{e^2}/\overline{v_n^2})$ to $(\overline{v^2}/\overline{v_n^2})$, i.e., by a factor of $(\overline{v^2}/\overline{e^2})$. This increase in signal-to-noise ratio may be read directly from the graph of the compression characteristic, and is unaffected by the identical treatment accorded signal and noise at the expander. Furthermore, noise received during the silent intervals, between speech bursts, is attenuated by the expander.

Under these circumstances it is appropriate to resolve companding improvement into the separate contributions of an increased signal to noise ratio for weak speech by the compressor, and a quieting of the circuit in the absence of speech by the expander. The introduction of an independent source of noise in the channel between the compressor and expander is the key to such behavior.

2. Instantaneous Companding of Unquantized Pulse Signals

Time division systems, employing unquantized pulse modulation (e.g., PAM) are admirably suited to the application of instantaneous companding to the individual pulse samples. Since each pulse is amplified to a degree which varies with its input amplitude, the compressor output is a sampled version of a distorted signal.

As in the syllabic case, the location of the noise source in the channel between the compressor and expander permits an improvement of the received signal-to-noise ratio for weak signals. Furthermore, the expander again assumes the separate and distinct task of suppressing channel noise in the absence of speech.

Unfortunately, quantitative expression of the companding improvement is not as simple as in the syllabic case. The response to instantaneous amplitudes much lower than the rms threshold signal (including zero) becomes important and the improvement factor may not (except in the special case of a linear compression characteristic) simply be read from a graph relating instantaneous values of v and e . Instead, one must employ the probability density of the signal in order properly to account for the distinctive treatment accorded individual pulse amplitudes in a complex signal.

3. Instantaneous Companding of Quantized Signals

Although the same physical devices which serve as an instantaneous compressor and expander in a PAM system may also be used in a PCM system, the functional description of companding improvement is different in the two applications. Whereas the compander is used to combat channel noise in a PAM system, encoded transmission permits a PCM system to assign this task to the devices which transmit and regenerate code pulses. Thus, assuming that error-free encoded transmission is realized, the quantized signal may be regarded as completely impervious to noise in the transmission medium. Quantization is required to permit such transmission. The sole purpose of the PCM compander is to reduce the quantizing impairment of the signal by converting uniform to effectively nonuniform quantization.

Although the expander continues to collaborate with the compressor in improving the quality of weak signals, it is now neither necessary nor possible for it to perform the separate function of quieting the circuit in the absence of speech. Indeed, apart from instrumentational difficulties which might arise, it is conceptually sound to transfer the PCM expander to the transmitting terminal, with expansion taking place subsequent to quantization but prior to encoding and transmission. Another interesting peculiarity of the PCM expander is the restriction of its operation, by quantization, to a finite number of discrete operating points on the continuous characteristic.

The use of companding to reduce the quantizing error which owes its very existence to, and is therefore a function of, the signal, is thus significantly different from the use of companding to reduce the effects of an independent source of noise in the transmission medium.

F. Applicability of the Present Analysis

Before we proceed to a detailed analysis, it is important to emphasize certain restrictive conditions required for the meaningful application of the results to be derived.

1. Signal Spectrum

A signal with a sufficiently complex spectrum, such as speech, is required to justify consideration of the total quantizing error power without regard to the detailed composition of the error spectrum. Although it is known that quantization of simple signals (e.g., sinusoids) results in discrete harmonics and modulation products deserving of individual attention,^{8, 9} Bennett has shown that the error spectrum for complex signals is sufficiently noise-like to justify analysis on a total power basis.^{2, 12}

2. Sampling Rate

The consistent comparison of signal power with the total quantizing error power, rather than with the fraction of the latter quantity appropriate to the signal band, might at first appear to impose serious limitations on the present analysis. Furthermore, the role of sampling has not been discussed explicitly. It is therefore important to note that the justification for this treatment, in the situation of actual interest, has also been given by Bennett.² We need only add the standard hypothesis¹⁻⁶ that the sampling rate chosen for a practical system would equal the

minimum acceptable rate (slightly in excess of twice the top signal frequency³) in order to invoke Bennett's results, which tell us that, for this sampling rate, the quantizing error power in the signal band and the total quantizing error power are identical.² Thus, sampling at the minimum rate is assumed throughout.

3. Number of Quantizing Steps

As already remarked, the present results are based on the assumption that N is not small, inasmuch as we assume a probability density which, although varying from step to step, remains effectively constant within each quantizing step; indeed the step sizes will be treated as differential quantities.

Experimental evidence^{6, 7, 10} (as well as the analysis to follow) argues against the consideration of fewer than five digits (i.e., $2^5 = 32$ quantizing steps) for high quality transmission of speech. Numerical estimates indicate that the present approximation should be reasonable for five or more digits per code group. These estimates are confirmed by the consistency of actual measurements of quantizing error power with calculations based on the same approximation (see Fig. 8 of reference 2 for 5, 6, and 7 digit data obtained with an input signal consisting of thermal noise instead of speech).

Further indication of the adequacy of this approximation is provided by the knowledge that Sheppard's corrections (see Section II-B) appear adequate even when (Δe) is not very small, for a probability density which (as is the case for speech¹⁸) approaches zero together with its derivatives at both ends of the (voltage) range under consideration.²⁴

Therefore, we are not presently concerned with the limitations imposed by this approximation.

4. Subjective Effects Beyond the Scope of the Present Analysis

We shall have occasion to study graphs depicting the signal to quantizing error power ratio as a function of signal power. Although these curves, and the equations they represent, will always be of interest for the case where even the weakest signal greatly exceeds the corresponding error power, there exists the possibility of rash extrapolation to the region where this inequality is reversed. Unfortunately, such extrapolation may have little or no meaning.* This is particularly clear when one considers that signals incapable of exciting at least the first quantizing step, in the absence of companding, will be absolutely incapable

* This is implicit in the deduction of Equation (6).

of transmission. Under these circumstances, companding may actually *resuscitate* a signal; the mathematical description of resuscitation (as anything short of infinite improvement) is clearly beyond the scope of the present analysis.

At the other extreme, it is probable that there exists a limit of error power suppression beyond which listeners will fail to recognize any further improvement. Our analysis will not be useful in describing this region of subjective saturation. Furthermore, it is possible that the subjective improvement afforded a listener by adding to the number of quantizing steps, or companding, may depend on the initial and final states, even before subjective saturation is reached. For example, it is entirely possible that the change from 5 to 6 digits per code group may provide a degree of improvement which appears different to the listener from that corresponding to the increase from 6 to 7 digits, although the present mathematical treatment does not recognize such a distinction.

II. EVALUATION OF MEAN SQUARE QUANTIZATION ERROR (σ)

A.* *Generalization of the Analysis of Panter and Dite*

The mean square error voltage, σ_j , associated with the quantization of voltages assigned to the j^{th} voltage interval, e_j , is adopted as the significant measure of the error introduced by quantization. If e_j is to represent any voltage, e , in the range

$$Q_j = \left[e_j - \frac{(\Delta e)_j}{2} \right] \leq e \leq \left[e_j + \frac{(\Delta e)_j}{2} \right] = R_j \quad (1)$$

then

$$\sigma_j = \int_{Q_j}^{R_j} (e - e_j)^2 P(e) de \quad (2)$$

where $(e - e_j)$ is the voltage error imparted to the sample amplitude by quantization and $P(e)$ is the probability density of the signal. The location of e_j at the center of the voltage range assigned to this level minimizes σ_j since we shall assume an effectively constant value of $P(e)$ within the confines of a single step.

If the value of $P(e)$ is approximated by the constant value $P(e_j)$ appropriate to e_j in (2), it follows that

$$\sigma_j = (\Delta e)_j^3 P(e_j)/12 \quad (3)$$

* This passage contains mathematical details which may be omitted, in a first reading, without loss of continuity.

The total mean square voltage error, σ , is equal to the sum of the mean square quantizing errors introduced at each level, so that,

$$\sigma = \sum_j \sigma_j = \frac{1}{12} \sum_j P(e_j) (\Delta e)_j^3 \quad (4a)$$

$$= \frac{1}{12} \sum_j (\Delta e)_j^2 [P(e_j) (\Delta e)_j] \quad (4b)$$

which may be rewritten as

$$\sigma \cong \frac{1}{12} \sum_j (\Delta e)_j^2 p_j \quad (4c)$$

since the discrete probability appropriate to the j^{th} step is given by

$$p_j = \int_{Q_j}^{R_j} P(e) de \cong [P(e_j) (\Delta e)_j] \quad (5)$$

Hence,

$$\sigma \cong \frac{1}{12} [(\Delta e)^2]_{AV} = \overline{(\Delta e)^2} / 12 \quad (6)$$

Thus, the total mean square error voltage is equal to one-twelfth the average of the square of the input voltage step size when the steps are sufficiently small (and therefore numerous) to justify the approximations employed in the deduction of (6). In applying (6), it is important to note that (4) implies that this is a *weighted* average over the steps traversed (or "excited") by the signal.

In the special case of uniform step size, substitution of $(\Delta e)_j = (\Delta e) = \text{const.}$ reduces (6) to the simple form

$$\sigma_0 = [\sigma]_{\Delta e = \text{const}} = (\Delta e)^2 / 12 \quad (7)$$

Equations (6) and (7) provide the basis for the qualitative interpretation of quantizing error power which has already been discussed in connection with Fig. 1.

The deduction of (6) from (4a) is implicit in the work of Panter and Dite.⁵ The absence of an explicit formulation of (6) therein* results from the direct application of the equivalent of (4b) to a specific problem involving a particular algebraic expression for $(\Delta e)_j$.

A prior, equivalent derivation of (7), based on a graphical representation of $(e - e_j)$ as a sawtooth error function for uniform quantization has been given by Bennett.² Although this derivation bypassed (6), Bennett has also analyzed compressed signals by means of an expression

* The present notation has been chosen to resemble that of Reference 5 in order to facilitate direct comparison by the reader.

[(1.6) of Reference 2] which is equivalent to (6), when the average is expressed as an integral over a continuous probability distribution and (Δe) is replaced by $(de/dv)(\Delta v)$, with $(\Delta v) = \text{const.}$ This form of (6) is the point of departure for the calculation in the Appendix.

B. Operational Significance of σ

Manipulation of (2) may be shown to result in the expression

$$\sigma = \sum_j \sigma_j = \sum_j e_j^2 p_j - \int e^2 P(e) de,$$

which is the difference between the mean square signal voltages following and preceding quantization. Hence σ is proportional to the difference between the quantized and unquantized signal powers. Since σ is intrinsically positive, the quantizing error power is *added* to the signal by quantization and is, in principle, measurable as the difference between two wattmeter readings.

In addition to providing an operational interpretation of the quantizing error power, the rewritten expression for σ reveals the equivalence of σ to the "Sheppard correction" to the grouped second-moment in statistics,²⁴⁻²⁷ where calculations are facilitated by grouping (i.e., uniform quantization) of numerical data. The reader who is interested in a more elaborate deduction of (7) from the Euler-Maclaurin summation formula, as well as discussions of the validity of (7), may therefore consult the statistical literature.

III. CHOICE OF COMPRESSION CHARACTERISTIC

A. Restriction to Logarithmic Compression

We shall consider the properties of the logarithmic type of compression characteristic,* defined by the equations†

$$v = \frac{V \log [1 + (\mu e/V)]}{\log (1 + \mu)}, \quad \text{for} \quad 0 \leq e \leq V \quad (8a)$$

and

$$v = \frac{-V \log [1 - (\mu e/V)]}{\log (1 + \mu)}, \quad \text{for} \quad -V \leq e \leq 0 \quad (8b)$$

* The author first encountered this characteristic in the work of Panter and Dite⁵ and the references thereto cited by C. P. Villars in an unpublished memorandum. He has since learned that such characteristics had been considered by W. R. Bennett as early as 1944 (unpublished), as well as by Holzwarth¹⁶ in 1949.

† Unless otherwise specified, natural logarithms will be used throughout.

In (8), v represents the output voltage corresponding to an input signal voltage e , and μ is a dimensionless parameter which determines the degree of compression.

Typical compression characteristics, corresponding to various choices of the compression parameter, μ , in (8a), are shown in Fig. 3. The logarithmic replot of Fig. 4 provides an expanded picture of small amplitude behavior, as well as evidence of the probable realizability of such characteristics.

Although restriction of attention to (8) may at first appear to impose serious limitations on the generality of the analysis, this impression is not confirmed by more careful scrutiny of the problem.

Perusal of Fig. 3 indicates that (8) generates a considerable variety of curves which meet the general requirements already enunciated in connection with Fig. 2. Thus, the constant factor, $V/\log(1 + \mu)$, has been chosen to satisfy the condition

$$[v]_{e=V} = V \quad (9)$$

Evidence of the significance of the μ -characteristics may be derived

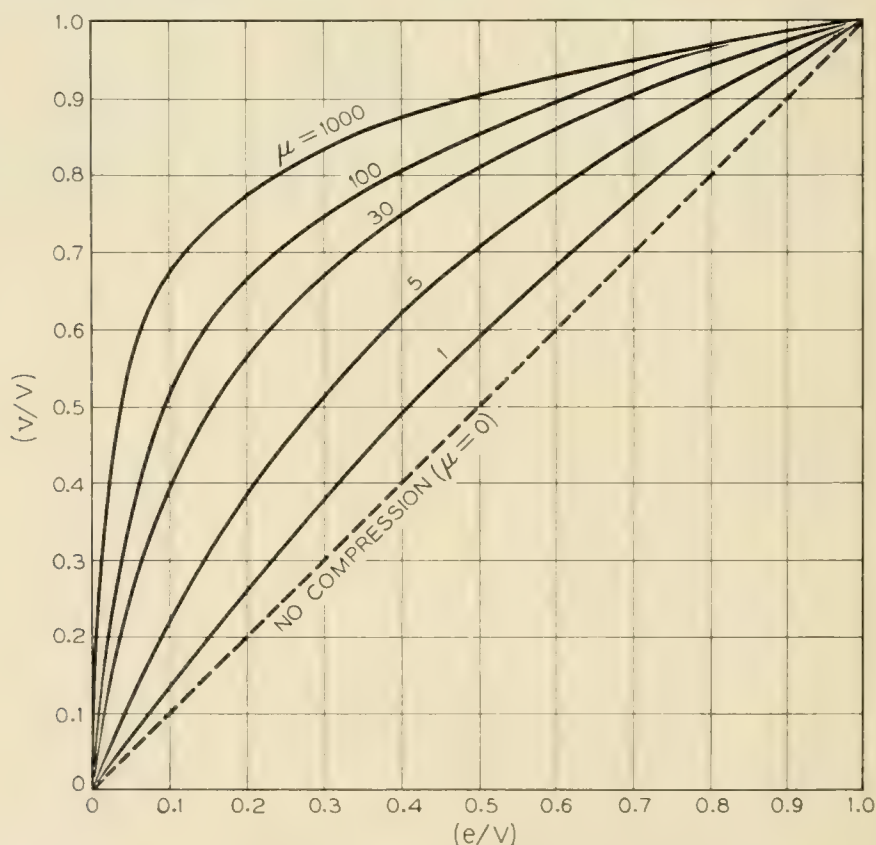


Fig. 3 — Typical logarithmic compression characteristics determined by equation (8a). The symmetrical negative portions, corresponding to equation (8b), are not shown.

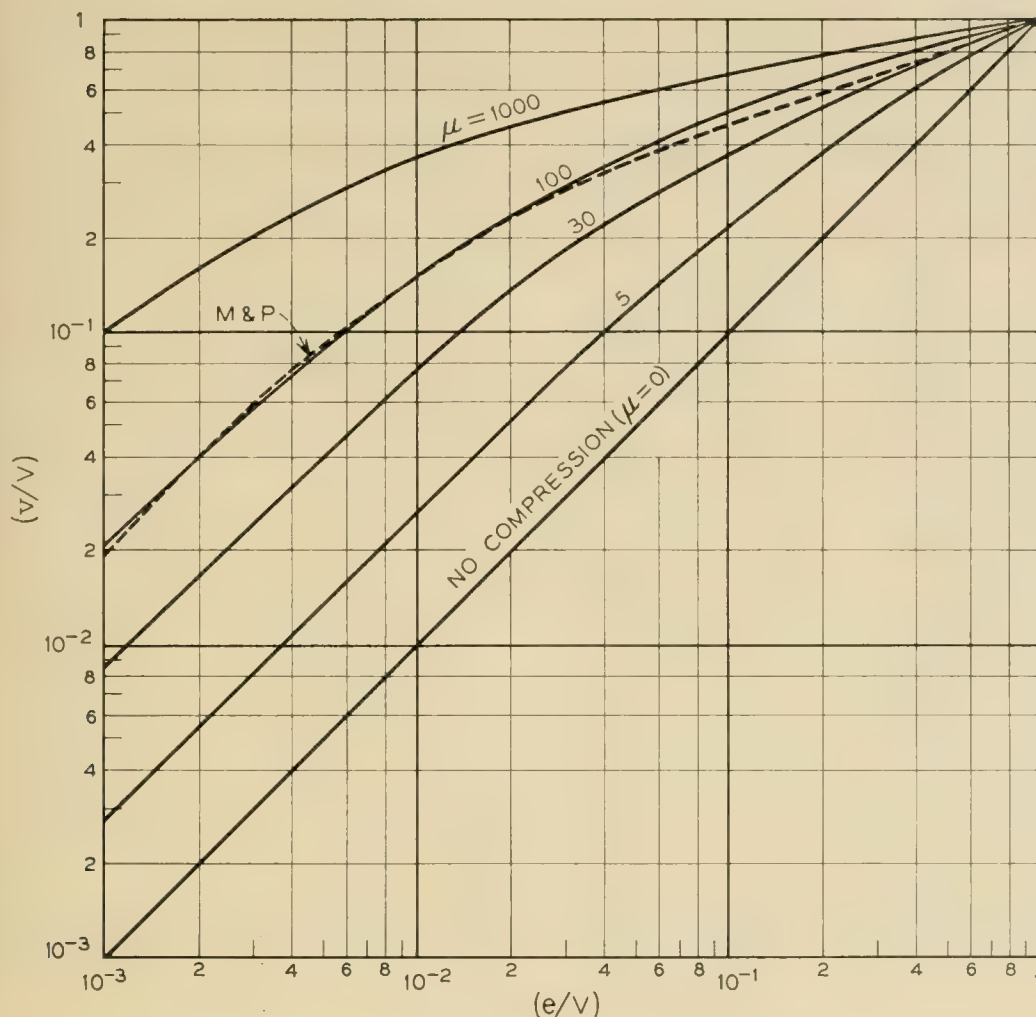


Fig. 4 — Logarithmic replot of compression curves shown in Fig. 3, to indicate detailed behavior for weak samples. The characteristic employed in the experiments of Meacham and Peterson⁶ (M & P) is also shown. Similarity between this characteristic and the $\mu = 100$ curve testifies to the probable realizability of these logarithmic characteristics.

from consideration of the ratio of step size to corresponding pulse amplitude, $(\Delta e/e)$, since this quantity is a measure of the maximum fractional quantizing error imposed on individual samples. Hence the relation,

$$(e/\Delta e) = [N/2 \log (1 + \mu)](1 + V/\mu e)^{-1}$$

[which follows from (12a)] has been plotted, for $\mu = 10, 100$, and 1000 , in Fig. 5. These curves reflect the fact that the sample to step size ratio reduces to the asymptotic forms:

$$(e/\Delta e) \rightarrow N/2 \log (1 + \mu) = \text{const} \quad \text{for} \quad (e/V) \gg \mu^{-1}$$

and

$$(e/\Delta e) \rightarrow [N/2 \log (1 + \mu)](\mu e/V) \quad \text{for} \quad (e/V) \ll \mu^{-1}$$

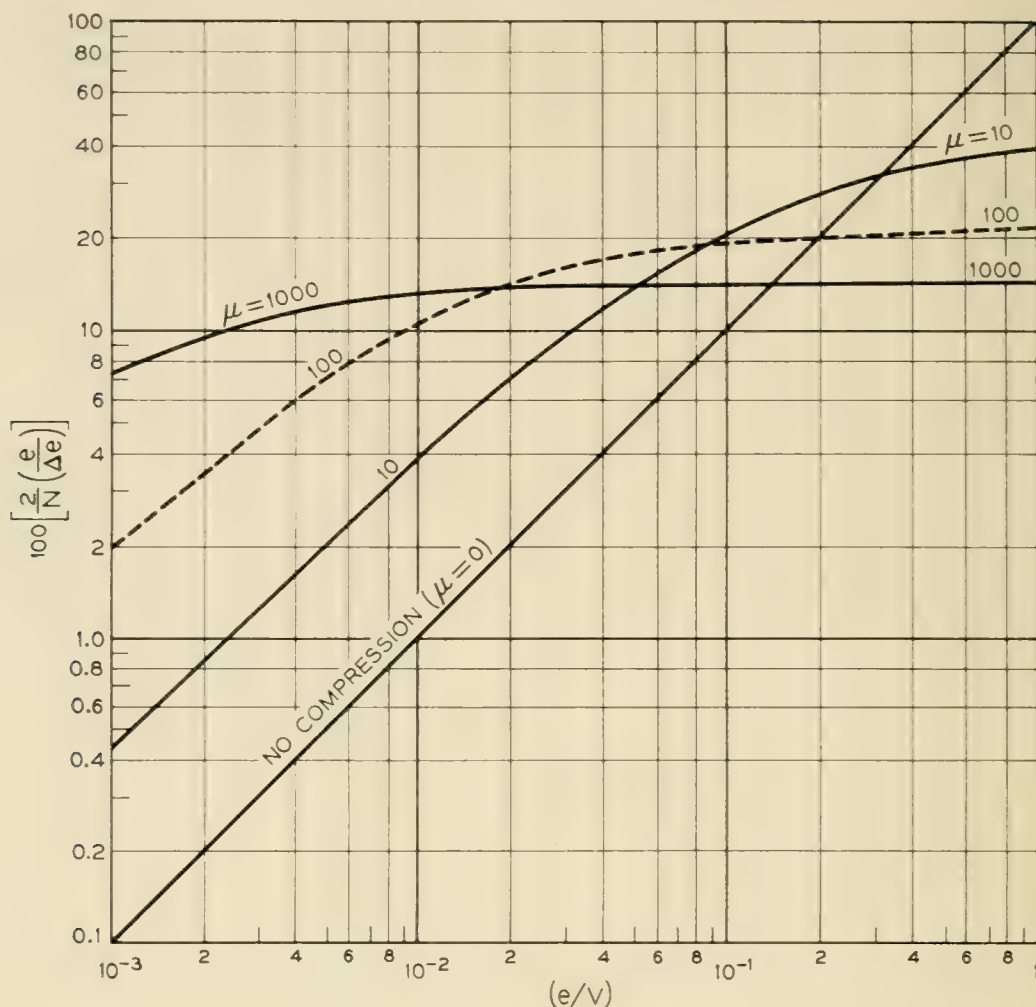


Fig. 5 — Pulse sample to step size ratios, as a function of relative sample amplitude, for various degrees of logarithmic companding (i.e., values of μ). The factor $(2/N)$ in the ordinate permits the curves to be drawn without reference to the total number of quantizing steps (N); the factor (100) is included to permit the ordinates directly to convey the proper order of magnitude for $(e/\Delta e)$, since present interest will be found to center about values of N for which $100(2/N) \sim 1$. As noted in the text, the ordinates, which constitute an index of the precision of quantization, approach constancy for $(e/V) \gg \mu^{-1}$, and vary linearly with abscissa for $(e/V) \ll \mu^{-1}$.

The essentially logarithmic behavior ($e/\Delta e \cong \text{const}$) for large pulse amplitudes is intuitively desirable since it implies an approach to the equitable reproduction of the entire distribution of amplitudes in a specified signal. Although existing experimental evidence indicates that the small amplitudes are not only most numerous,¹⁸ but also most significant for the *intelligibility*²¹⁻²³ of speech at constant volume, the absence of comparable evidence on the properties of *naturalness* makes it plausible to consider only those compression characteristics which give promise of providing the same, acceptably small, upper limit on the fractional quantizing error for pulse samples of all sizes.

For sufficiently small input pulses, $(e/\Delta e)$ becomes proportional to e , as a result of the linearity of the logarithmic function in (8) for small arguments. In view of our professed preference for logarithmic behavior, with $(e/\Delta e) \cong \text{const.}$, it is important to emphasize that the transition to linearity is not peculiar to (8), but is rather an example of the linearity to be expected of any suitably behaved (i.e., continuous, single-valued, with $(dv/de)_{e=0} > 1$) odd compression function, $v(e)$, capable of power series expansion, in the vicinity of the origin. In (8) this transition to linearity takes place where (e/V) is comparable to μ^{-1} . The extension of the region where $(e/\Delta e) \cong \text{const.}$ to lower and lower pulse amplitudes requires an increase in μ , and a concomitant reduction of the $(e/\Delta e)$ ratio for strong pulses.

Further evidence of the significance of the parameter μ may be deduced by evaluating the ratio of the largest to the smallest step size from the asymptotic expressions for $(e/\Delta e)$. Thus we find

$$\frac{(\Delta e)_{e=V}}{(\Delta e)_{e=0}} \rightarrow \mu \quad \text{for} \quad \mu \gg 1$$

which is a special form of the more general relation⁵

$$\frac{(\Delta e)_{e=V}}{(\Delta e)_{e=0}} = \frac{(dv/de)_{e=0}}{(dv/de)_{e=V}} = \mu + 1$$

which follows from our standard approximation of

$$(de/dv) \cong (\Delta v/\Delta e)$$

with $\Delta v = \text{const.}$

B. Comparison with Other Companders

An upper bound for companding improvement, which permits the quantitative evaluation of the penalty incurred (if any) through the restriction to logarithmic companding, is established in the Appendix. Comparison of the results to be derived from (8) with this upper bound will reveal that nonlogarithmic characteristics, which provide somewhat more companding improvement at certain volumes, are apt to prove too specialized for the common application to a broad volume range envisioned herein. The μ -characteristics do not suffer from this deficiency since the equitable treatment of large samples, which we have hitherto associated with an "intuitive naturalness conjecture," will be seen to tend to equalize the treatment of all signal volumes.

Finally, it will develop that (8), when applied to (6), has the added merit of calculational simplicity.

IV.* THE CALCULATION OF QUANTIZING ERROR

A. *Logarithmic Companding in the Absence of "DC Bias"*

As previously noted, we consider the effect of *uniformly* quantizing a *compressed* signal. If we designate the uniform output voltage step size by (Δv) , then

$$(\Delta v) = \frac{2V}{N} \quad (10)$$

since the full voltage range between $-V$ and $+V$, of extent $2V$, is to be divided into N equal steps. For a number of levels, N , which is sufficiently large to justify the substitution of the differentials dv and de for the step sizes Δv and Δe , differentiation of (8a) yields

$$\frac{(\Delta v)}{V} = k \left[\frac{1}{1 + (\mu e/V)} \right] \frac{\mu(\Delta e)}{V} \quad (11)$$

where $k = 1/\log(1 + \mu)$.

Combining (10) with (11) and the counterpart of the latter in the domain of (8b), we find

$$\Delta e = \alpha(V + \mu e) \quad \text{for} \quad 0 \leq e \leq V \quad (12a)$$

and

$$\Delta e = \alpha(V - \mu e) \quad \text{for} \quad -V \leq e \leq 0 \quad (12b)$$

where

$$\alpha = 2 \log(1 + \mu)/\mu N \quad (13)$$

Substitution of (12) into (6) yields

$$\sigma = (\alpha^2/12)[V^2 + \mu^2 \bar{e}^2 + 2\mu V \overline{|e|}] \quad (14)$$

where the quantity $\overline{|e|}$ is introduced by the difference in sign in (12a) and (12b). For ordinary compandor applications, we may write

$$\overline{|e|} = 2 \int_0^V e P(e) de \quad (15)$$

since the symmetry of the input signal provides that $P(-e) = P(e)$ and $\bar{e} = 0$.

* This passage contains mathematical details which may be omitted, in a first reading, without loss of continuity.

It is convenient to define the quantization error voltage ratio,

$$D = \frac{\text{RMS Error Voltage}}{\text{RMS Input Signal Voltage}} = (\sigma/\bar{e}^2)^{\frac{1}{2}} \quad (16)$$

which takes the form

$$D = \log(1 + \mu)[1 + (C/\mu)^2 + 2AC/\mu]^{\frac{1}{2}}/\sqrt{3N} \quad (17)$$

when we define the quantities

$$A = \overline{|e|}/\sqrt{\bar{e}^2} = \frac{\text{Average Absolute Input Signal Voltage}}{\text{RMS Input Signal Voltage}} \quad (18)$$

and

$$C = V/\sqrt{\bar{e}^2} = \frac{\text{Compressor Overload Voltage}}{\text{RMS Input Signal Voltage}} \quad (19)$$

The simple linear proportionality of Δe to $(V \pm \mu e)$ results from the properties of the logarithmic function in differentiation. Other, seemingly more simple compression equations, when differentiated, yield much more complicated and unwieldy expressions for Δe . The value of this simplicity is evident in the absence, from (14), of moments of e higher than the second.

If we set $A = 0$, (17) reduces to one deduced by Panter and Dite⁵; their analysis erroneously associated A with $\bar{e} = 0$ rather than with $\overline{|e|}$, as a result of their tacit assumption that (12a) and (12b) are identical. They also imposed the restriction of considering only that class of input signals having peak values coincident with the compandor overload voltage, by defining V as the peak value of the signal in specifying C . The definition of C in terms of the independent properties of both signal ($\bar{e}^2$) and compandor (V) is then converted into one based solely on the properties of the signal. This interpretation leads to conclusions quite different from those to be presented here.

B. *Logarithmic Companding in the Presence of "DC Bias"*

It has heretofore been assumed that the input signal is symmetrically disposed about the zero voltage level since it may be expected that $\bar{e} = 0$ for speech. Although this is a standard assumption, subsequent discussion will disclose that it is probable, in actual practice, for the average value of the input signal to be introduced at a point other than the origin of the compression characteristic. In terms of Fig. 1, the signal is

presented to the array of quantizing steps with its quiescent value displaced by an amount e_0 from the center of the voltage interval ($-V$ to $+V$).

Such an effect, regardless of its origin, may formally be described by considering the composite input voltage

$$E = e + e_0 \quad (20)$$

where e is the previously considered symmetrical speech signal and e_0 is the superimposed constant voltage.

Substitution of E for e in (8) and (12) yields

$$\sigma_E = (\alpha^2/12)[V^2 + \mu^2 \overline{E^2} + 2\mu V \overline{|E|}] \quad (21)$$

where the subscript E is introduced to distinguish this result from (14). Note that the value $[e]_{E=0} = -e_0$ now separates the domain of applicability of (8a) and (12a) from that of (8b) and (12b), so that (15) is replaced by

$$\overline{|E|} = \int_{-V}^{-e_0} (-E)P(e) de + \int_{-e_0}^V EP(e) de \quad (22)$$

which reduces to

$$\overline{|E|} = \overline{|e|} + 2e_0 \int_0^{e_0} P(e) de - 2 \int_0^{e_0} eP(e) de \quad (23)$$

Since $\bar{e} = 0$, and $e_0 = \text{const.}$, we also find

$$\overline{E^2} = \bar{e}^2 + e_0^2 \quad (24)$$

C. Application to Speech as Represented by a Negative Exponential Distribution of Amplitudes

It is necessary to assume an explicit function for $P(e)$ in (15) and (23) before applying the general results which have thus far been deduced. We shall assume, as a simple but adequate first approximation, that the distribution of amplitudes in speech at constant volume¹⁸ may be represented by

$$P(e) = G \exp(-\lambda e) \quad \text{for} \quad e \geq 0 \quad (25)$$

where $P(-e) = P(e)$, $G = \lambda/2$, and $\lambda^2 = 2/\bar{e}^2$. The values of G and λ

follow from the standard relations

$$\int_{-\infty}^{\infty} P(e) de = 1$$

and

$$\int_{-\infty}^{\infty} e^2 P(e) de = \bar{e}^2$$

When applied to (15) and (18), with the upper limit in (15) replaced by ∞ with negligible error, (25) implies that

$$A = |\overline{e}| / \sqrt{\bar{e}^2} = 1/\sqrt{2} = 0.707 \quad (26)$$

Hence, (17) will be replaced, for numerical calculations, by the relation

$$\sqrt{3}ND = \log (1 + \mu)[1 + (C/\mu)^2 + \sqrt{2}C/\mu]^{\frac{1}{2}} \quad (27)$$

The corresponding substitution of (25) into (23) yields, for the case of $e_0 \neq 0$,

$$|\overline{E}| = e_0 + (\bar{e}^2/2)^{\frac{1}{2}} \exp(-\sqrt{2}C/B) \quad (28)$$

where we have introduced the "bias parameter,"

$$B = V/e_0 \quad (29)$$

When (28) is combined with (13), (21), and (24), we find, after some algebraic manipulation, that

$$\begin{aligned} \sqrt{3}ND_E &= \log (1 + \mu) \\ &\cdot [1 + (C/\mu)^2(1 + \mu/B)^2 + (\sqrt{2}C/\mu) \exp(-\sqrt{2}C/B)]^{\frac{1}{2}} \end{aligned} \quad (30)$$

where $D_E^2 = (\sigma_E/\bar{e}^2)$. It is to be noted that D_E has been defined in terms of the ratio of σ_E to $\bar{e}^2$ rather than $\bar{E}^2$, so that

$$D_E^2 = \frac{\text{Mean Square Error Voltage}}{\text{Mean Square Speech Voltage}} \quad (31a)$$

$$= \frac{\text{Average Error Power}}{\text{Average Speech Power}} \quad (31b)$$

Examination of (30) reveals that it has the required property of reducing to (27) for $e_0 = 0$, i.e., for $B \rightarrow \infty$. Furthermore (27) and (30) indicate that $D_E \geq D$ so that the addition of a dc component increases the quantizing error power when companding is used. The existence of

such an impairment may easily be understood in terms of the physical interpretation of (6), as discussed in connection with Fig. 1.

Equations (27) and (30) also reveal that the penalty inflicted by a finite e_0 is largely determined by the ratio (μ/B) . If $(\mu/B) \ll 1$, the presence of e_0 will be unimportant. At the other extreme, if $(\mu/B) \gg 1$, $(1 + \mu/B)^2 \rightarrow (\mu/B)^2$ and

$$\sqrt{3}ND_E \rightarrow \log(1 + \mu) \cdot [1 + (C/B)^2 + (\sqrt{2}C/\mu) \exp(-\sqrt{2}C/B)]^{\frac{1}{2}} \quad (32)$$

which proves to be relatively insensitive to changes in μ for the values of μ , C and B considered herein. In this case B largely usurps the algebraic role previously assigned to μ in (27).

D. Uniform Quantization: $\mu = 0$

The mean square quantization voltage error in the absence of companding, corresponding to direct, uniform quantization of the input signal, follows immediately from (7) and (10) since $\Delta v = \Delta e$ under these conditions. Thus

$$\sigma_0 = (\Delta v)^2/12 = V^2/3N^2$$

whence

$$D_0 = (\sigma_0/\bar{e}^2)^{\frac{1}{2}} = C/\sqrt{3}N \quad (33)$$

This inverse proportionality of D_0 and N is well known.^{2, 3, 5}

Equation (33) may also be deduced by letting μ approach zero in the expressions for D and D_E , since (8) implies that v approaches e as μ approaches zero. The fact that $D_0 = (D_E)_{\mu \rightarrow 0}$ reveals that, in the absence of companding, the addition of e_0 does not change the quantizing error power. This conclusion was anticipated in the discussion of Fig. 1.

V. DISCUSSION OF GENERAL RESULTS

Since the nature of a companded signal depends on a rather large number of variables, it is appropriate to consider their respective roles in general terms before discussing detailed system requirements. This general discussion will, however, emphasize those particular modes of operation which are suggested by existing proposals for application of PCM.^{1, 3, 6} Thus, we shall consider common channel companding of

speech in time division multiplex systems which employ binary number encoding.

A. Quantitative Description of Conventional Operation ($e_0 = 0$)

1. Number of Quantizing Steps (N)

The number of steps, N , is related to the choice of code. For a binary code, the relation takes the form, $N = 2^n$, where n is the number of binary digits per code group.

In the present discussion it will usually be convenient to regard n and N as fixed in order to permit comparison of various companding characteristics under the same coding conditions. Since both D and D_0 are inversely proportional to N , the quotient $(D/D_0)^2$, which is equal to the ratio of the quantizing error power in the presence of companding to that in the absence of companding, is independent of N . Consequently, as will be evident in the discussion of (37), the relative diminution of quantizing error (in db) afforded by companding is also independent of N .*

However, the value of N will determine the signal to quantizing error power ratio to which the companding improvement is to be added. Thus the number of digits per code group required for a particular application will ultimately be determined by the value of N which, in combination with suitable companding, will suffice to produce an acceptably low value of quantizing error power in relation to signal power.

2. Compandor Overload Voltage (V)

The compandor overload voltage, V , will be determined by the full load power objectives for the proposed system. Specifically, V will be equal to the amplitude of the sinusoidal voltage corresponding to "full modulation."

3. Relative Signal Power (C)

The quantity $C = V/(\bar{e}^2)^{\frac{1}{2}}$ will, for a given value of V , be determined by the rms signal voltage $(\bar{e}^2)^{\frac{1}{2}}$.

The range of C values appropriate to a given system will therefore reflect the distribution of volumes to be encountered. In fact, the signal

* It must of course be understood that this independence requires a value of N sufficiently large to justify the approximations involved in the deduction of (27) and (33).

power in db *below* that corresponding to a full load sine wave is simply

$$10 \log_{10} \left[\frac{V^2/2}{e^2} \right] = 10 \log_{10} [C^2/2] = 20 \log_{10} C - 3 \text{ db} \quad (34)$$

It must be emphasized that C varies *inversely* with the rms signal voltage, so that weak signals are characterized by large values of C and strong signals by small values of C .

4. Average Absolute Signal Amplitude ($\overline{|e|}$)

The probability density of the signal manifests itself in the value assigned to the average absolute signal parameter, $A = \overline{|e|} / \sqrt{e^2}$. This

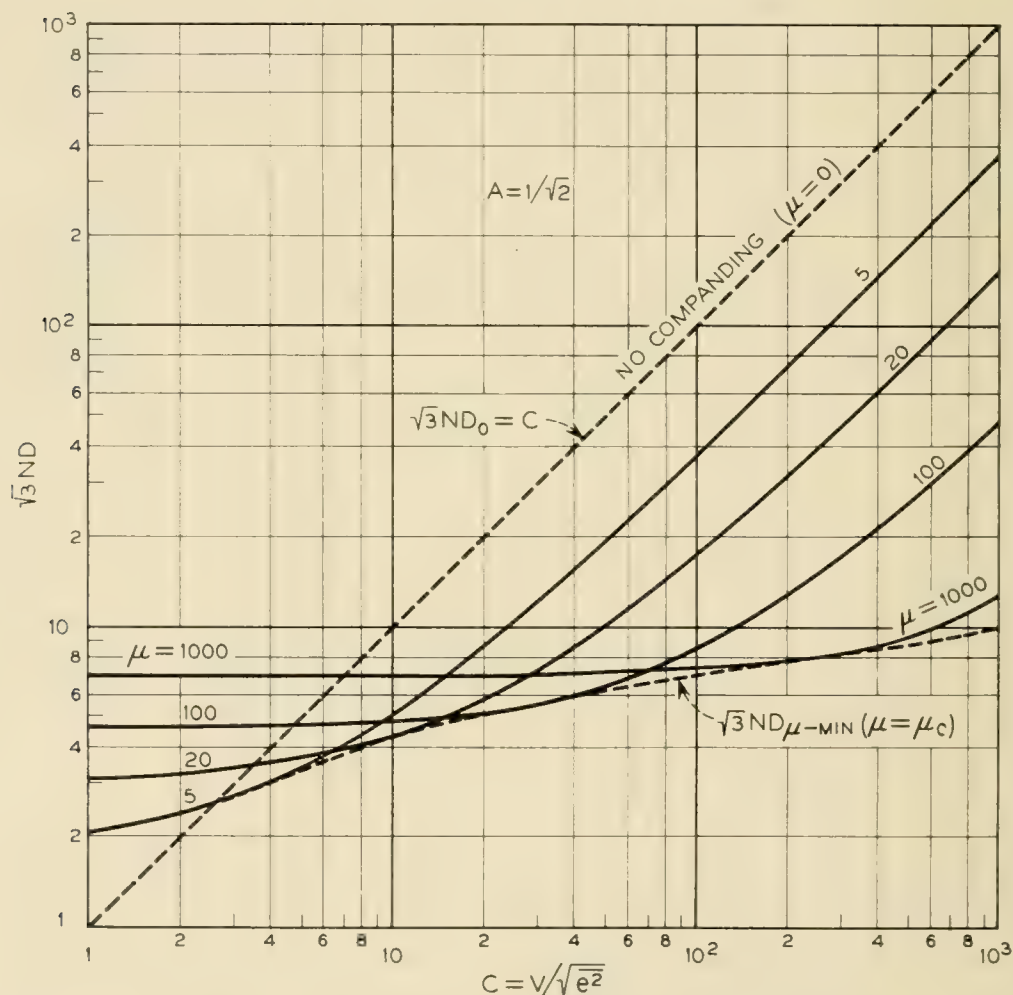


Fig. 6 — Variation of the rms error to signal voltage ratio (D) with relative signal strength, $C = V/\sqrt{e^2}$, as given by equation (27) for various degrees of logarithmic companding.

quantity is a constant determined by the statistical properties of the class of signals being studied.

With the present choice of an exponential distribution of amplitudes to represent speech, [see (25)], we have seen that A takes on the value $1/\sqrt{2} = 0.707$. It develops that A is not very sensitive to the choice of $P(e)$, as may be judged by the values $\sqrt{2/\pi} = 0.798$, and $\sqrt{3}/2 = 0.866$ which would replace 0.707 if (25) were replaced by Gaussian and rectangular distributions, respectively. The value $A = 1/\sqrt{2}$ will be used in all numerical calculations; changes in the value of A to describe other classes of signals (e.g., the aforementioned Gaussian or rectangular distributions) will change the plotted results by no more than a fraction of a decibel.

5. Degree of Compression (μ)

From the foregoing it is clear that the essence of the compandor's behavior is embodied in the one remaining variable which appears in (8) and (17): the compression parameter μ .

The significance of μ has already received preliminary attention in connection with Figs. 3 to 5. Fig. 6, where comparison of behavior at constant N is facilitated by the choice of $\sqrt{3}ND$ as ordinate, exhibits the behavior of the ratio D as a function of C at constant μ . It will be observed that the curves in Fig. 6 do not extend below their common tangent which is labeled $\sqrt{3}ND_{\mu-\text{MIN}}$. The significance of this lower bound may be discussed in terms of Fig. 7 and the hypothetical ensemble of compandors to which we now direct our attention.

B. Optimum Compandor Ensemble

Consider the artificial situation in which our communication system includes an ensemble of instantaneous compandors, the members of which correspond to different values of μ in (8). Since companding improvement varies with signal strength, we permit ourselves the luxury of measuring the volume (i.e., C) of the input signal in order to assign the optimum degree of companding compatible with (8), to each individual signal. The compandor assigned to a signal is characterized by that particular value of the compression parameter, $\mu = \mu_c$, which is required to minimize D for a particular value of C . This critical compression parameter may be calculated from the requirement that

$$[\partial D / \partial \mu]_{A, C = \text{const}} = 0 \quad (35)$$

which yields

$$SC^2 + \mu_c A(S - 1)C - \mu_c^2 = 0 \quad (36)^*$$

where $S = [(1 + \mu_c) \log(1 + \mu_c)/\mu_c] - 1$, when applied to (17).

The graph of (36) in Fig. 7 may be used to determine numerical values of μ_c without repeated recourse to the equation.

The curve labeled $\sqrt{3}ND_{\mu-\text{MIN}}$ in Fig. 6 was determined by substituting values of μ_c , obtained from Fig. 7, into (27). Each curve in Fig. 6

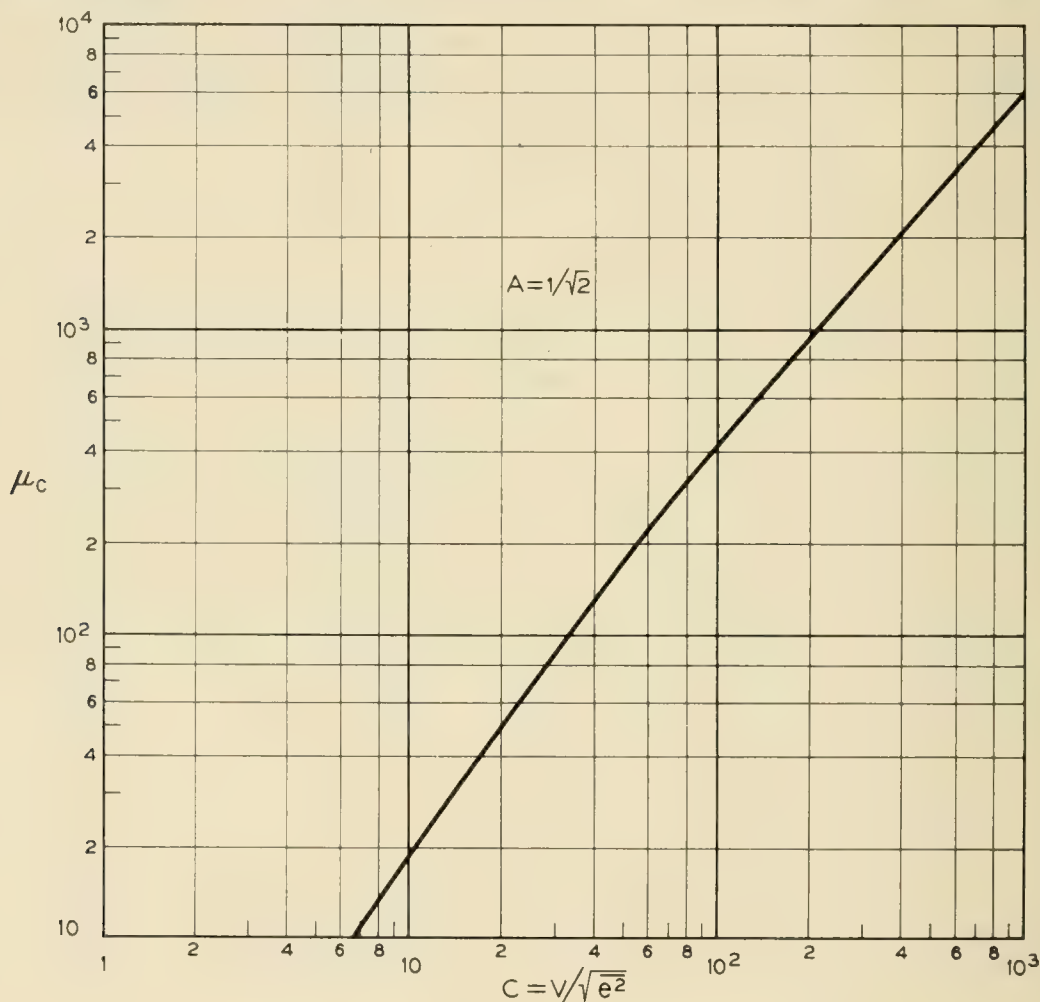


Fig. 7 — Critical compression parameters, μ_c , required to minimize the quantizing error power as a function of relative signal strength, as determined by equation (36). Each point on the curve defines a compandor in the optimum compandor ensemble. It must be understood that such an ensemble provides the best performance consistent with equation (8) rather than the absolute minimum quantizing error discussed in the Appendix.

* A similar equation, with $A = 0$ corresponding to the previously noted erroneous identification of A with $\bar{e}$ rather than $|\bar{e}|$, has been deduced by Panter and Dite.⁵ Their definition of C as a “crest factor” changes the significance of what we have called $D_{\mu-\text{MIN}}$ and does not lead to the ensemble interpretation of μ_c .

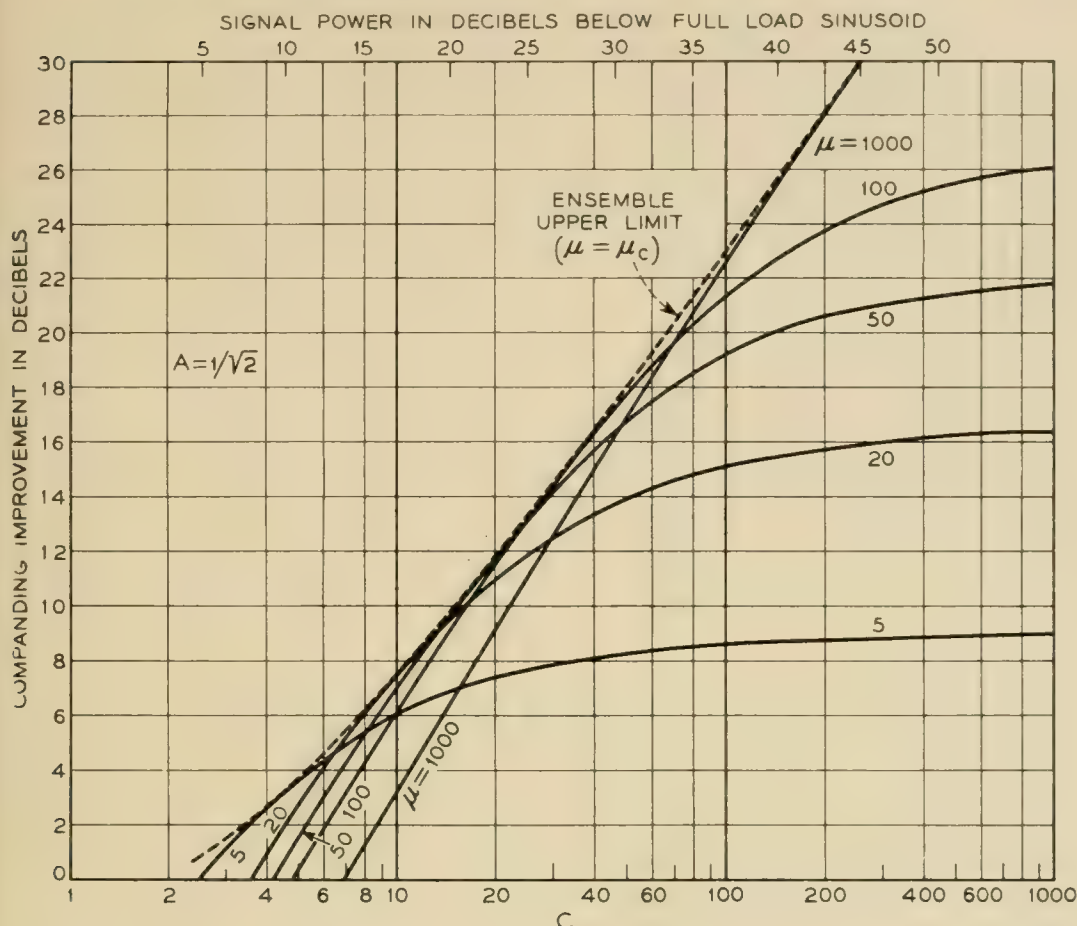


Fig. 8 — Companding improvement (in db), as calculated from equation (37), for various values of μ . The saturated improvement for weak signals (relatively constant ordinate for large C values) is identical with the asymptotic behavior for weak signals which is predicted by Fig. 9.

is tangent to this lower bound at the single value of C which corresponds to $\mu = \mu_c$.

In conventional systems, a single common channel compandor,⁶ characterized by a single value of μ , is substituted for the optimum ensemble. Although $D_{\mu-\text{MIN}}$ is then attainable at only one value of C , it is instructive to compare each value of D with the corresponding value of $D_{\mu-\text{MIN}}$. Indeed, consideration of the optimum ensemble has, in one sense, reduced the problem of choosing an appropriate μ for a given application to the choice of that particular value of C at which equality of D and $D_{\mu-\text{MIN}}$ is desired.

In Fig. 6, the line representing performance in the absence of companding corresponds to (33) for D_0 . D_0 and $D_{\mu-\text{MIN}}$ are seen to be similar for strong signals (low values of C). Furthermore, it is important to note that D_0 does not constitute an upper bound for D ; thus the companding

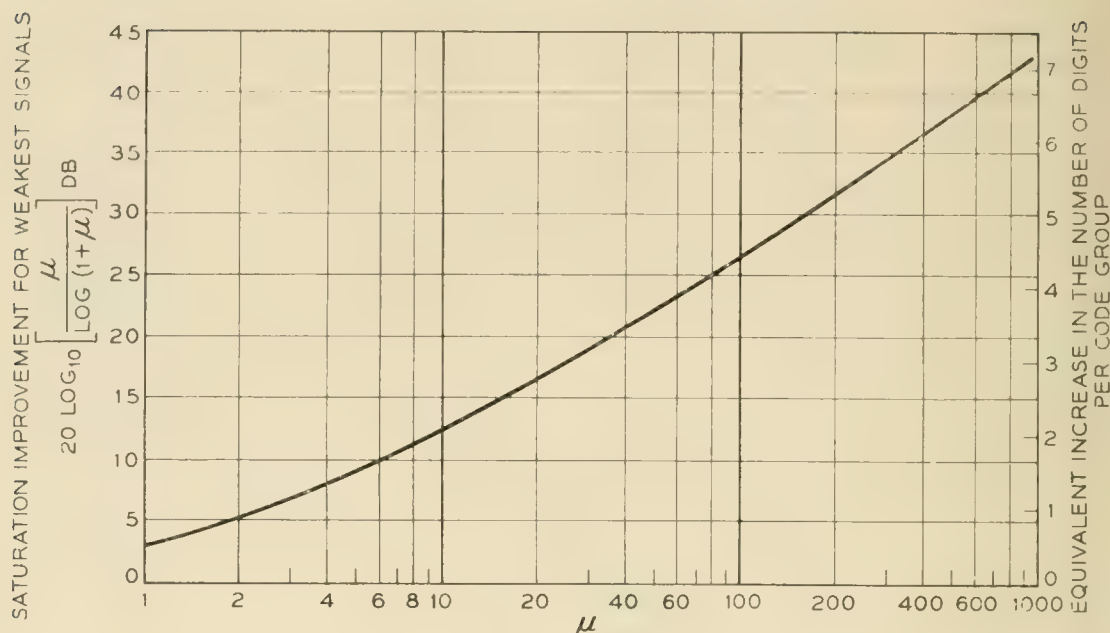


Fig. 9 — Saturated companding improvement for weakest signals as a function of the degree of logarithmic compression (μ). Given a value of μ , the corresponding ordinate represents the reduction of quantizing error power (in db) for signals so weak that their peaks satisfy the relation $(e/V) \ll \mu^{-1}$. Thus weaker and weaker signals are required to exploit the added improvement which follows from an increase in μ . The auxiliary ordinate scale on the right exhibits the increase in the number of digits per code group which may be shown to be equivalent to the linear compression which gives rise to the saturated improvement of weak signals.

improvement of weak signals (large C) is obtained at the price of impairment of strong signals (small C).

C. Companding Improvement for $e_0 = 0$

From the definition of D in (16) it follows that

$$\begin{aligned} \left[\frac{D_0}{D} \right]^2 &= \frac{\text{Mean Square Error Voltage Without Companding}}{\text{Mean Square Error Voltage With Companding}} \\ &= \frac{\text{Average Quantizing Error Power Without Companding}}{\text{Average Quantizing Error Power With Companding}} \end{aligned}$$

whence, the improvement (in db) due to companding may be evaluated from the expression

$$10 \log_{10} (D_0/D)^2 = 20 \log_{10} (D_0/D) \text{ db} \quad (37)$$

This quantity is plotted against C for various values of μ in Fig. 8. Substitution of $D_{\mu-\text{MIN}}$ for D in (37) yields the optimum ensemble limiting curve. Just as in Fig. 6, each curve is tangent to this limiting curve for a value of C corresponding to coincidence of μ and μ_c .

The abscissa has been furnished with an additional scale, based on (34), showing the signal power in db below that of a full load sine wave. Thus the curves in Fig. 8 embody the information required for the choice of a compandor characteristic in a form which is suitable for practical applications. It is clear that the effect of companding varies considerably with the volume of the signal and that strong signals are actually impaired. For each value of μ , there is a value of C (in the weak signal, or large C range) beyond which the db improvement is essentially independent of C . The threshold value of C corresponding to such saturation increases with increasing μ . This saturated improvement for weak signals reflects the linearity of compression for $(e/V) \ll \mu^{-1}$. In this region, uniform quantization prevails so that the ratio of the uniform step sizes before and after companding (given by the linear amplification factor) may be used to deduce a constant companding improvement for the

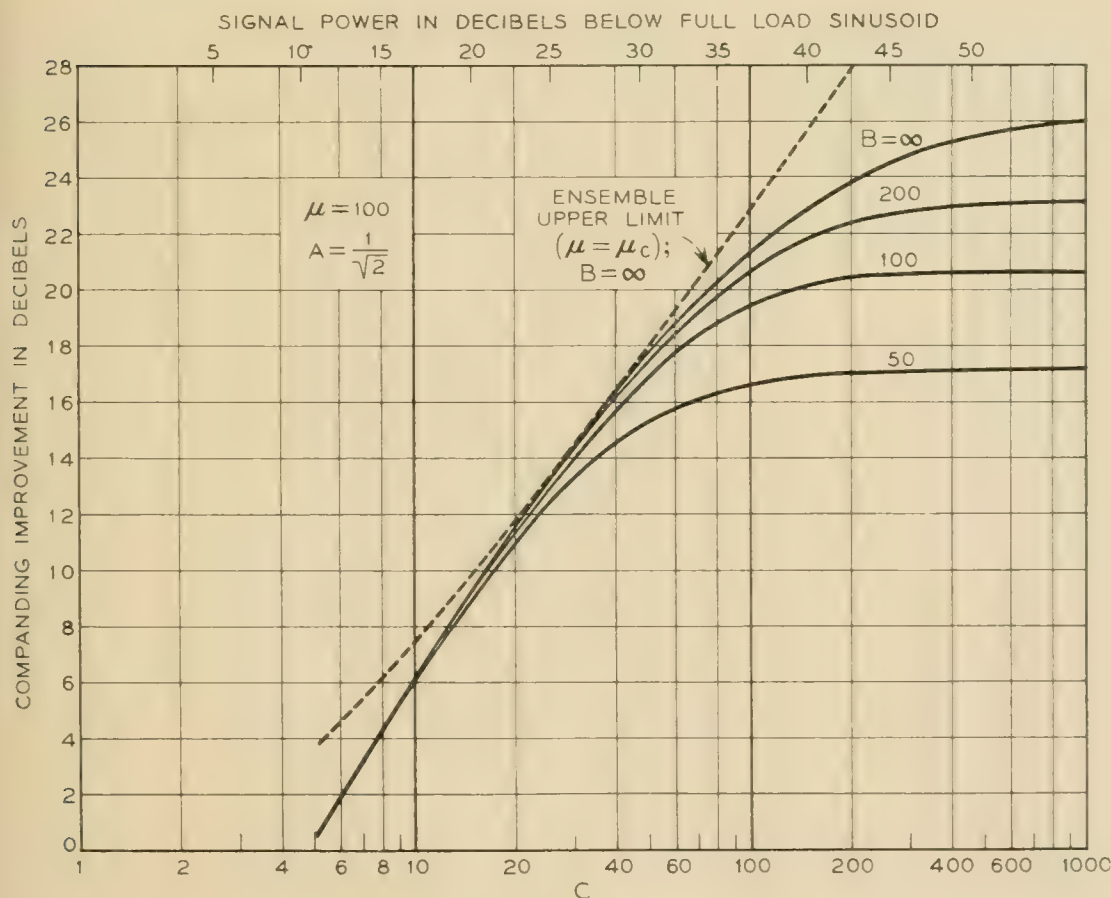


Fig. 10 — Modification of companding improvement, for $\mu = 100$, resulting from the shift of the quiescent value of the signal by an amount e_0 from the center of the quantized voltage range (see Fig. 1). The relative size of the “dc component” is given by the parameter $B = V/e_0$ as defined in equation (29). The algebraic usurpation of the role of μ by B , for $B \ll \mu$, results in the striking similarity of the weak signal behavior of the curves for $B = 50$ and 100 in Figs. 10 to 13.

weakest signals, regardless of the shape of the characteristic in the non-linear region. Equation (8) implies

$$(v/e)_{e \rightarrow 0} = (\Delta v / \Delta e)_{e \rightarrow 0} = \mu / \log(1 + \mu)$$

so that for each value of μ , the constant companding improvement for the weakest signals is

$$20 \log_{10} [\mu / \log(1 + \mu)] \text{ db}$$

which is plotted in Fig. 9. Fig. 8 supplements Fig. 9 in revealing the actual volumes required for the realization of this weak signal saturation, as well as the detailed behavior for stronger signals.

D. Companding Improvement for $e_0 \neq 0$

Since we will usually regard a nonzero value of e_0 as an undesirable

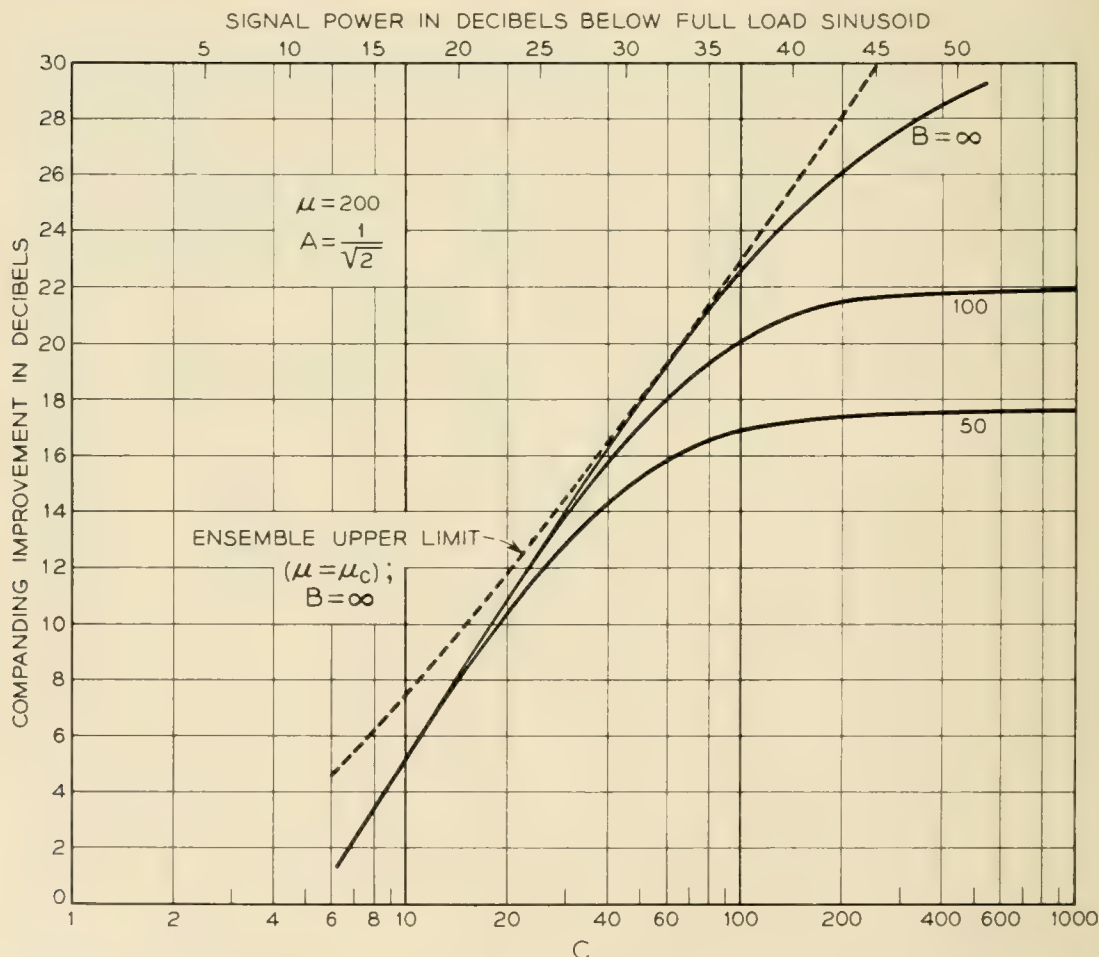


Fig. 11 — The effect of a “dc component” on companding improvement for $\mu = 200$. For further details, see the caption of Fig. 10.

perturbation, we wish to study the modification of companding improvement produced by the introduction of a finite value of $B = V/e_0$, i.e., substitution of (30) for (27), when N, V, C, A , and μ remain unchanged.

In Figs. 10 to 13 we have replotted the companding improvement curves shown in Fig. 8 for $\mu = 100, 200, 500$, and $1,000$, respectively. These curves correspond to $B = \infty$. The difference between these curves and those for finite values of B in Figs. 10 to 13 is the impairment (in db) inflicted by the presence of e_0 . This impairment may be appreciable for weak signals (large C). As already noted in connection with (32) the impairment is not severe for $(\mu/B) \ll 1$. Furthermore, the appropriation of the algebraic role of μ by B , when $B \ll \mu$, which was previously noted in (32), manifests itself in the striking similarity of the weak signal behavior of all the curves for $B = 50$ and 100 in Figs. 10 to 13.

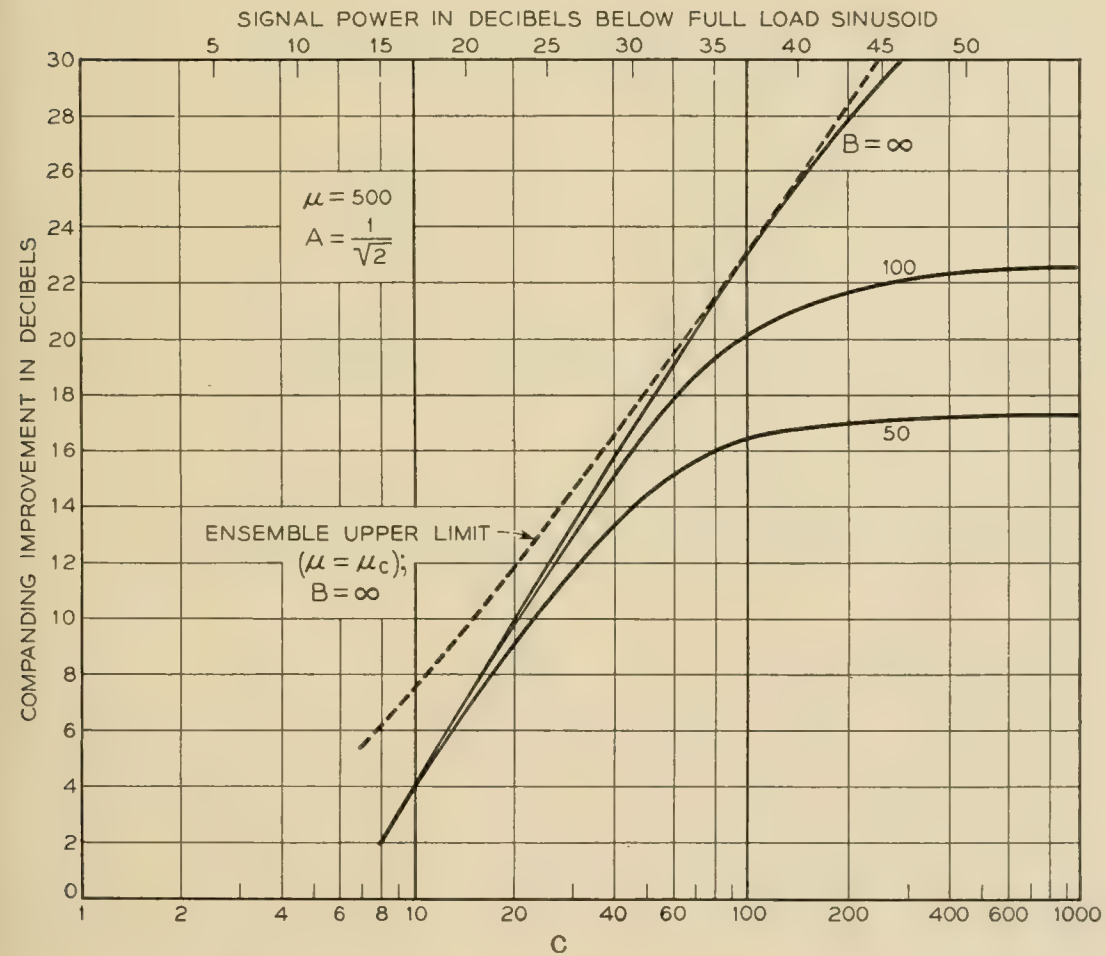


Fig. 12 — The effect of a “dc component” on companding improvement for $\mu = 500$. For further details, see the caption of Fig. 10.

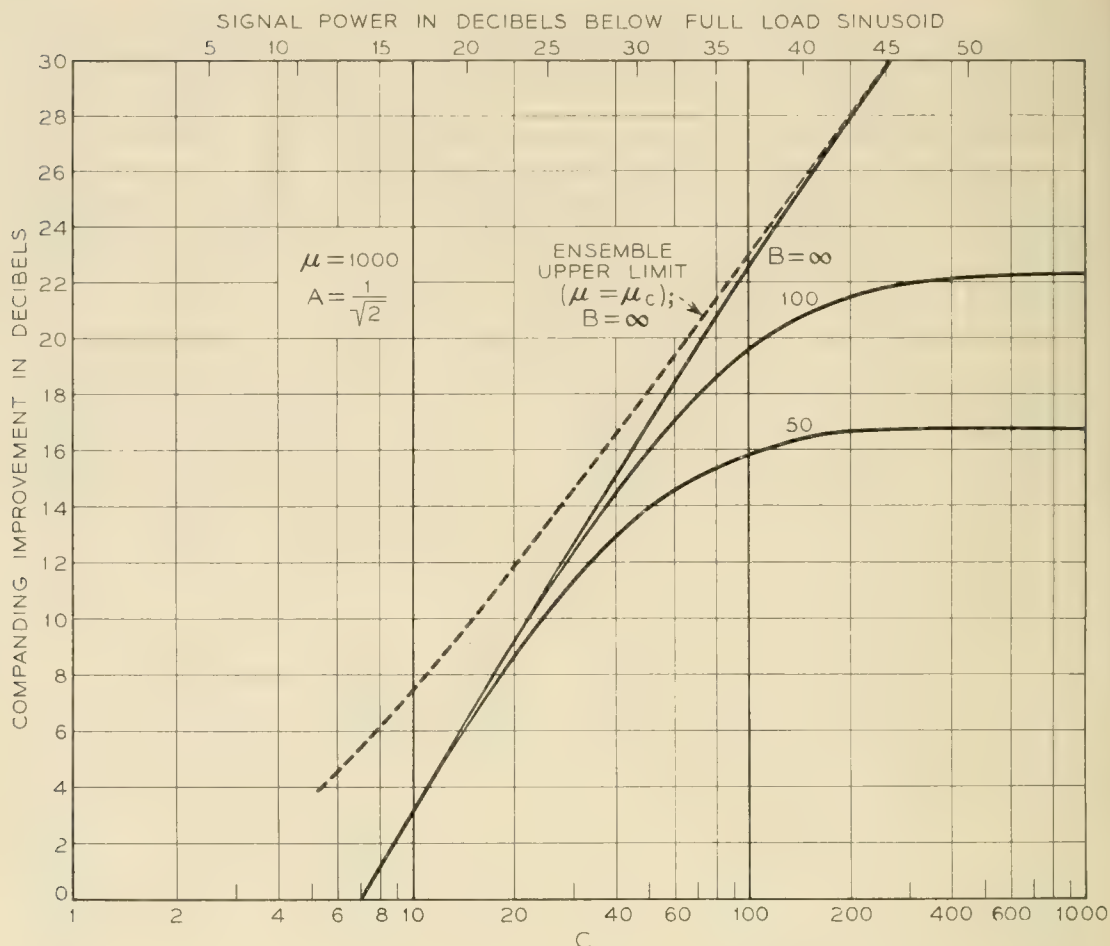


Fig. 13 — The effect of a “dc component” on companding improvement for $\mu = 1,000$. For further details, see the caption of Fig. 10.

VI. APPLICATION OF RESULTS TO A HYPOTHETICAL PCM SYSTEM

Consider the application of these results to the planning of a typical, albeit hypothetical, communication system.

A. Speech Volumes

Suppose it is desired to transmit signals covering a 40 db power range, with the strongest and weakest speech volumes each separated by 20 db from the average anticipated signal power at the compressor input.*

The strongest signal power is then used to determine the value of the compandor overload voltage, V . In this case a value of V corresponding to a full load sine wave 10 db above the loudest signal, [see (34)], appears adequate.* Although this choice may at first appear arbitrary,

* These values are sufficiently close to those cited as representative by Feldman and Bennett, in connection with Fig. 2 of Reference 11, to be considered quite realistic.

the value of 10 db is probably no more than a few db removed from the value which would be chosen in any efficient application of PCM to quality telephony. It results from the need to balance the requirement of a value of V sufficiently high to avoid intolerable clipping of the peaks of the loudest signals against the obvious advantage of reducing the quantizing step size by minimizing the voltage range to be quantized. We have neglected clipping in our calculations since it was assumed that the significant peaks of the loudest signals should not exceed V for quality telephony. Existing information on clipped speech²¹⁻²³ and one digit PCM¹⁰ indicates that the clipping impairment we seek to avoid is largely one of loss of naturalness rather than reduction in intelligibility. The choice of a maximum volume 10 db below a sinusoid of amplitude V implies that speech peaks 13 db, [see (34)], above the maximum rms signal voltage are being ignored, which appears reasonable in the light of available experimental evidence.^{18, 22}

It follows from these assumptions that the average and weakest signals are respectively 30 db and 50 db below full sinusoidal modulation.

B. Choice of Compression Characteristic

1. Ideal Behavior for Speech

If we adopt the aim of achieving the smallest over-all departure from the ensemble limit of improvement, it seems reasonable to choose that compandor in the optimum ensemble which corresponds to average speech ($C \cong 45$). This requirement, in conjunction with Fig. 7, establishes a lower bound of about 150 for μ . The significance of this choice may be clarified by reference to Fig. 14, which depicts departures from the optimum ensemble limit of improvement, resulting from restriction to a single value of μ for all volumes.

The corresponding upper bound will be determined by the alternative of furnishing optimum improvement to the weakest signals ($C \cong 450$) in spite of the concomitant impairment of loud speech. Reference to Fig. 7 then dictates a choice of μ in the vicinity of 2,500. From Fig. 6 it is clear that this value implies that D is essentially constant and independent of C throughout the range of interest. Appreciably larger values of μ would actually lead to the undesirable extreme of $D > D_{\mu-\text{MIN}}$ for all signals under consideration.

We therefore conclude that attention may profitably be confined to the interval $150 \lesssim \mu \lesssim 2,500$, the magnitude of which is adequately conveyed by the simple expression

$$100 \lesssim \mu \lesssim 1,000 \quad (38)$$

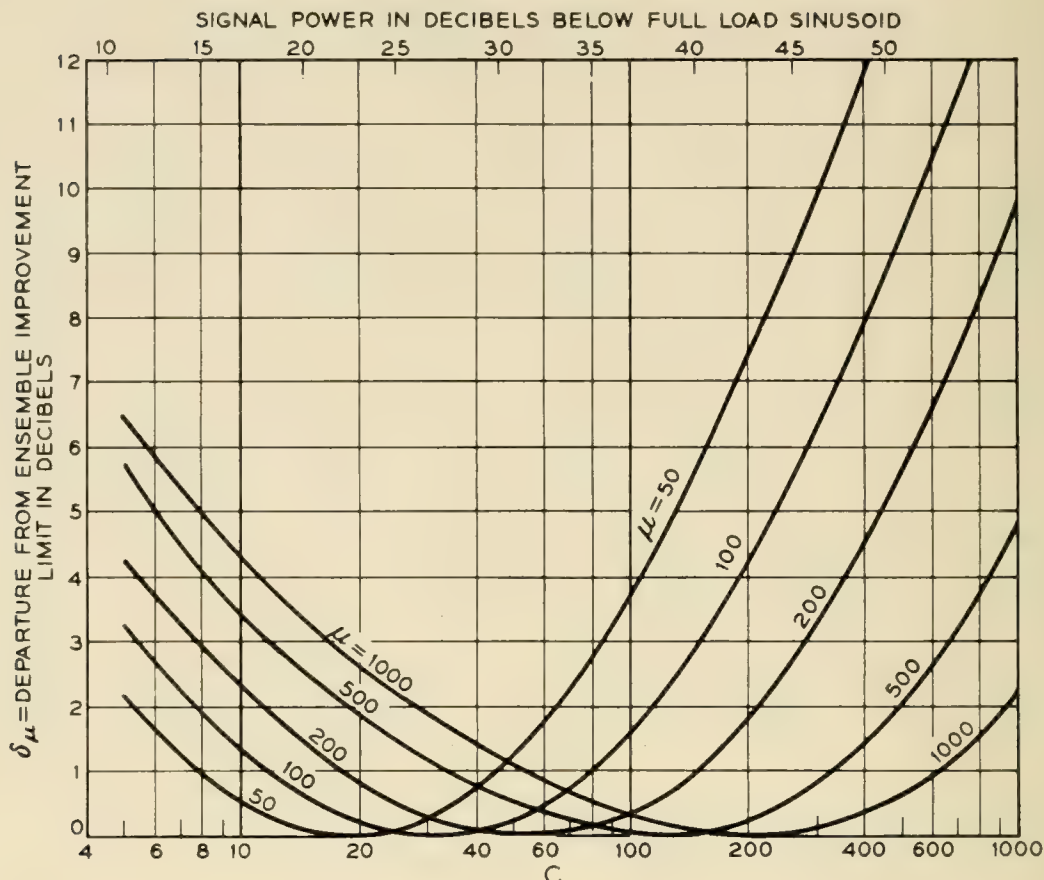


Fig. 14 — Improvement selectivity, i.e., departures from the ensemble upper limit of improvement due to the use of a single value of μ rather than the ensemble of μ_c values. The minima at $\delta_\mu = 0$ locate the signals (C) for which μ and μ_c coincide. All curves correspond to the case where the “dc component”, e_0 , is zero.

Lest it appear that this range is so broad as to offer very little practical guidance, it should be noted that (38) defines a rather narrow range of characteristics in Figs. 3 and 4. The assumption that this range may be realized in practice appears reasonable in view of the similarity to the characteristic actually used by Meacham and Peterson, which is shown in Fig. 4.

2. Practical Limitations on Companding Improvement

(a) *Mismatch Between Zero Levels of Signal and Compandor.* Although the present discussion has hitherto been confined to ideal compandor action, it lends itself quite naturally to the analysis of a significant departure from ideal behavior which may be expected to result from the use of an instantaneous compandor on a common channel basis in time division multiplex systems.

It will probably be impractical to balance the channel gating circuits (required to provide sequential connection of individual channels to a

single compressor)^{1, 6} sufficiently to guarantee exact coincidence of the average input signal ($\bar{e} = 0$) in each channel and the center of the $-V$ to $+V$ voltage range ($e = 0$) presented by the compressor. Thus the input, e , would appear to the compressor in the form $E = e + e_0$. The consequences of the appearance of the undesirable constant term, e_0 , may be inferred from study of Figs. 10 to 13 and (30).

We shall assume that, owing to the present state of gating technology, $B = V/e_0$ may reasonably be expected to assume values in the range $100 \lesssim B \lesssim 1,000$.

For companding corresponding to $100 \lesssim \mu \lesssim 1,000$, Figs. 10 to 13 indicate that, if B can be confined to the vicinity of 1,000, the departure from the ideal behavior corresponding to $B = \infty$ will be virtually negligible.

However, should it prove necessary to work with $B \cong 100$, it is clear from Figs. 10 to 13 that the companding improvement for weak signals would be relatively independent of μ in the interval $100 \lesssim \mu \lesssim 1,000$ (with a saturation value of about 20.5–22.5 db).^{*} In this event, compression to a degree greater than that represented by $\mu = 100$ would provide less improvement for strong and average speech without the compensation of significantly greater improvement for weak signals. Reduction of μ below 100 would not be fruitful since the sensitivity of companding improvement to changes in μ is restored for values satisfying the condition of $(\mu/B) = (\mu/100) < 1$.

The significance of the values $e_0 \sim V/1,000$ and $V/100$ may perhaps better be appreciated in terms of a comparison of e_0 with the weakest signals under consideration. Since $(B/C) = \sqrt{\bar{e}^2}/e_0$, a signal to dc bias power ratio may be calculated, in db, from the expression $20 \log_{10} (B/C)$. For the weakest signals under consideration ($C \sim 400$), the values $B = 1,000$ and 100 correspond respectively to $(\sqrt{\bar{e}^2}/e_0) = 2.5$ and 0.25, or to signal to dc bias power ratios of +8 db and -12 db. Thus, for the hypothetical system now under study, the value of e_0 becomes significant (roughly) when it exceeds the weakest rms signal.

Actually e_0 would be expected to vary with time for a given channel and to vary from channel to channel at any instant. On the assumption that $|e_0| = V/100$ (i.e., $B = 100$) will constitute the upper bound of such variations, the companding improvement corresponding to a particular value of μ must now be specified in terms of the region between the $B = \infty$ and $B = 100$ curves in Figs. 10 to 13, rather than by reference to a single value of B and its corresponding curve. Since the lower

^{*} This corresponds to the behavior of D_E for $(\mu/B) \gg 1$ which was noted in the discussion of (30). In this connection, see the discussion of Fig. 19.

bounds of all these regions (see Figs. 10 to 13) are approximately coincident (for $100 \lesssim \mu \lesssim 1,000$), the advantage of increasing μ substantially beyond 100 will depend largely on the expectation of encountering values of $c_0 \rightarrow (V/1,000)$ with sufficient frequency in the various channels served by the common compressor.

These arguments may of course be applied, with suitable modifications depending on the range of C , μ , and B values requiring attention, to any effect capable of formal description in terms of an effective dc bias superimposed on the signal input to the compressor.

(b) *Background Noise Level.* It does not seem reasonable to strive for an increase of the signal to quantizing error power ratio substantially beyond that value which is subjectively equivalent to the anticipated ratio of signal to background noise from other sources.

Since the quantizing error power depends on the number of digits per code group, the comparison of quantizing error power and noise power is reserved for subsequent discussion of the required number of quantizing steps. It will be noted that the comparison must remain somewhat speculative in the absence of a determination of the subjective equivalence of quantizing error power and noise.

C. Choice of the Number of Digits Per Code Group

1. Ideal Behavior for Speech

As previously remarked, the number of quantizing steps will determine the ratio of signal to quantizing error power to which the companding improvement is to be added. Since the quantizing error power is inversely proportional to $N^2 = 2^{2n}$, this power will be reduced by 6 db for each additional digit. Comparison of this 6 db per digit improvement with the roughly 24 to 35 db improvement corresponding to weak signals in Fig. 8 (for $100 \lesssim \mu \lesssim 1,000$) reveals that, *for such signals, companding is equivalent to the addition of four to six digits per code group, i.e., to an increase in the number of quantizing steps by a factor between $2^4 = 16$ and $2^6 = 64$.* This equivalence is portrayed in Fig. 9. Our failure to realize a companding improvement of about 43 db as predicted for $\mu = 1,000$ in Fig. 9 may be traced to the fact that the weakest signals now under consideration are not sufficiently weak to be confined to the linear region ($c/V \ll \mu^{-1}$) of the $\mu = 1,000$ characteristic. This is reflected in the unsaturated improvement exhibited in Fig. 8 for the weakest signals when $\mu = 1,000$.

Although it is clearly preferable to suppress quantizing error power by companding rather than by increasing the number of quantizing

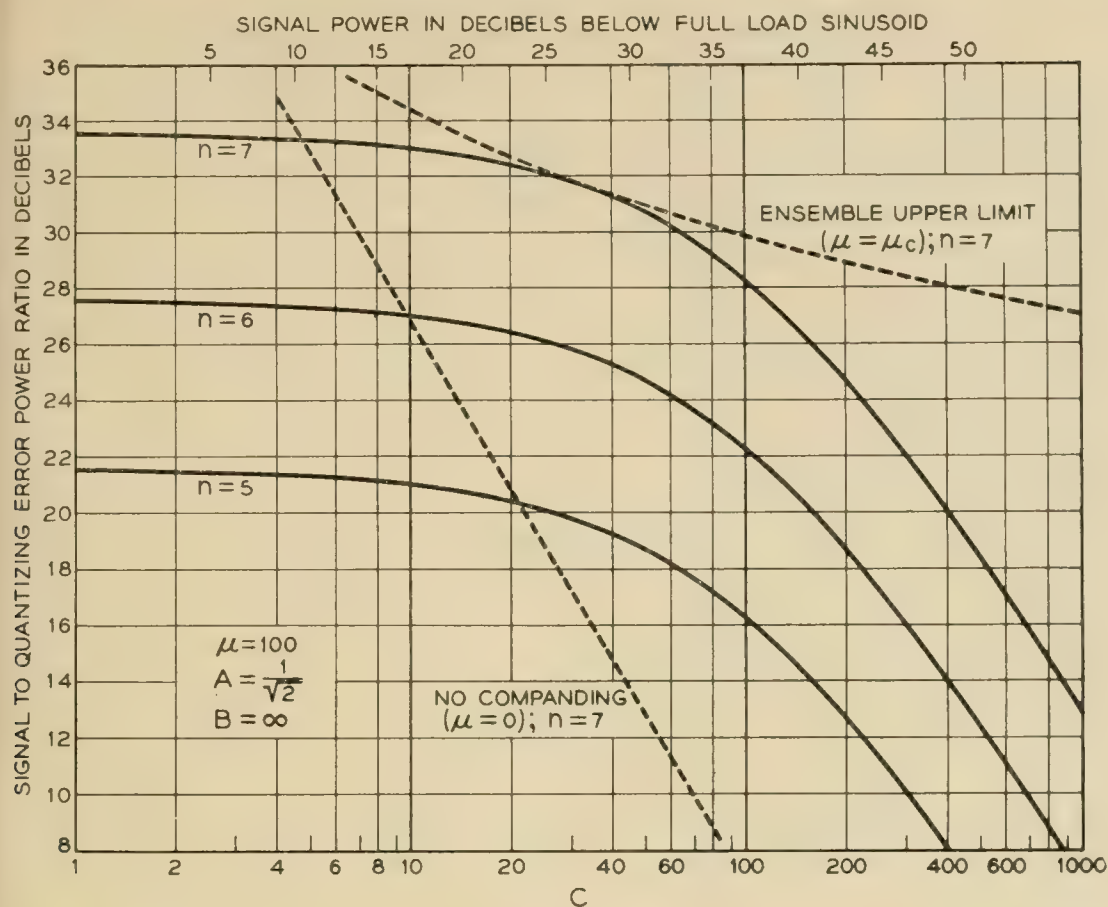


Fig. 15 — Signal to quantizing error power ratios (calculated, in db, from equation (39)) as a function of relative signal power for companding corresponding to $\mu = 100$. Curves are shown for $n = 5, 6$, and 7 digits per code group. For comparison, the results for seven digits in the absence of companding ($\mu = 0$) as well as for the ensemble upper limit ($\mu = \mu_c$) are included. $B = \infty$ throughout.

steps, it is apparent that the upper limit of companding improvement will set a lower limit on the number of digits required for satisfactory operation.

Once again we begin with the consideration of pure speech signals. The expression

$$\begin{aligned}
 -10 \log_{10} (D^2) &= -20 \log_{10} D \\
 &= \text{Signal to Quantizing Error Power Ratio in db}
 \end{aligned}
 \tag{39}$$

has been plotted against C in Figs. 15 to 18 for $\mu = 100, 200, 500$, and $1,000$ respectively. In each case the behavior for $5, 6$, and 7 digits is compared with the extremes of $\mu = 0$ (no companding) and $\mu = \mu_c$ (ensemble upper limit) for 7 digits.

It must be conceded at the outset that experimental work is required to formulate standards for quantizing error power similar to those

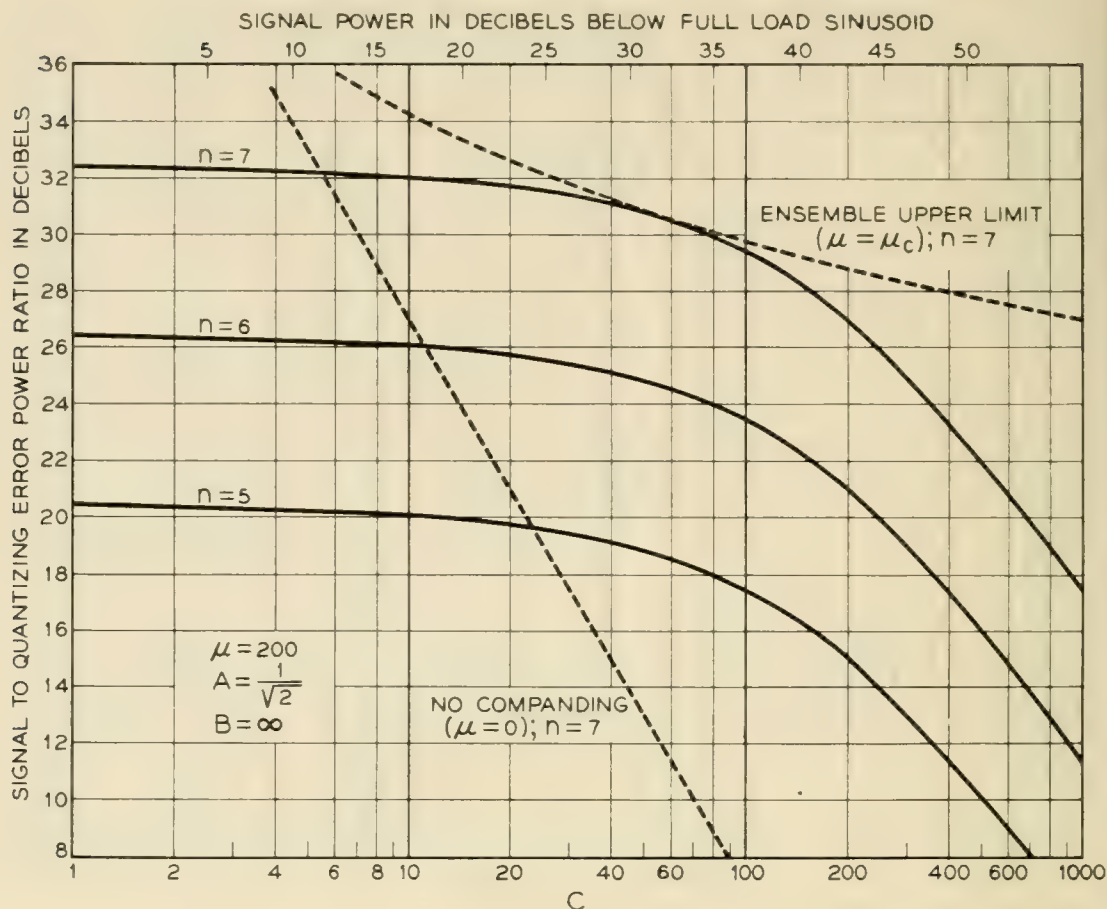


Fig. 16 — Signal to quantizing error power ratios (in db) as a function of relative signal power for companding corresponding to $\mu = 200$. Symbols have the same significance as in Fig. 15. $B = \infty$ throughout.

which have been established for conventional noise and distortion. If these were available, graphs such as those in Figs. 15 to 18 could be used to select the proper number of digits to be used with various degrees of compression. In the absence of such information we shall complete this illustrative study by adopting a signal to quantizing error power ratio of at least 20 db as a tentative standard of adequate performance at all volumes.*

Figs. 15 and 16 show that seven digits (i.e., $2^7 = 128$ tapered quantizing steps) and $\mu \cong 150$ will meet this objective. Furthermore Figs. 17 and 18 indicate that six digits ($2^6 = 64$ tapered steps) would suffice provided (38) is replaced by the more stringent limitation,

$$500 \lesssim \mu \lesssim 1,000 \quad (40)$$

* This value does not appear unreasonable, as a first approximation, in terms of experience with noise and harmonic distortion.

2. Practical Limitations

(a) *Mismatch Between Zero Levels of Signal and Compandor.* From the previous discussion of the effect of e_0 on the choice of μ , it is clear that, if B can be confined to the vicinity of 1,000, the analysis of the required number of digits in the absence of instability ($B = \infty$) may be applied.

On the other hand, behavior for $B = 100$ may be judged from the plot of signal to quantizing error power ratio versus signal power for $\mu = 100$ and 1,000 (with seven digits) shown in Fig. 19. Since this ratio now fails to exceed about 16 db for the weakest signals of interest, we conclude that an increase to eight digits ($2^8 = 256$ tapered steps), with a concomitant 6 db improvement for all signals, is required to meet our 20 db objective. These curves also illustrate the previously noted meager improvement for weak speech which accompanies the increase from $\mu = 100$ to 1,000 when $B = 100$. Actually, an optimum solution is attained for an intermediate value of μ , but the advantage is too small to be of interest (see Figs. 10 to 13).

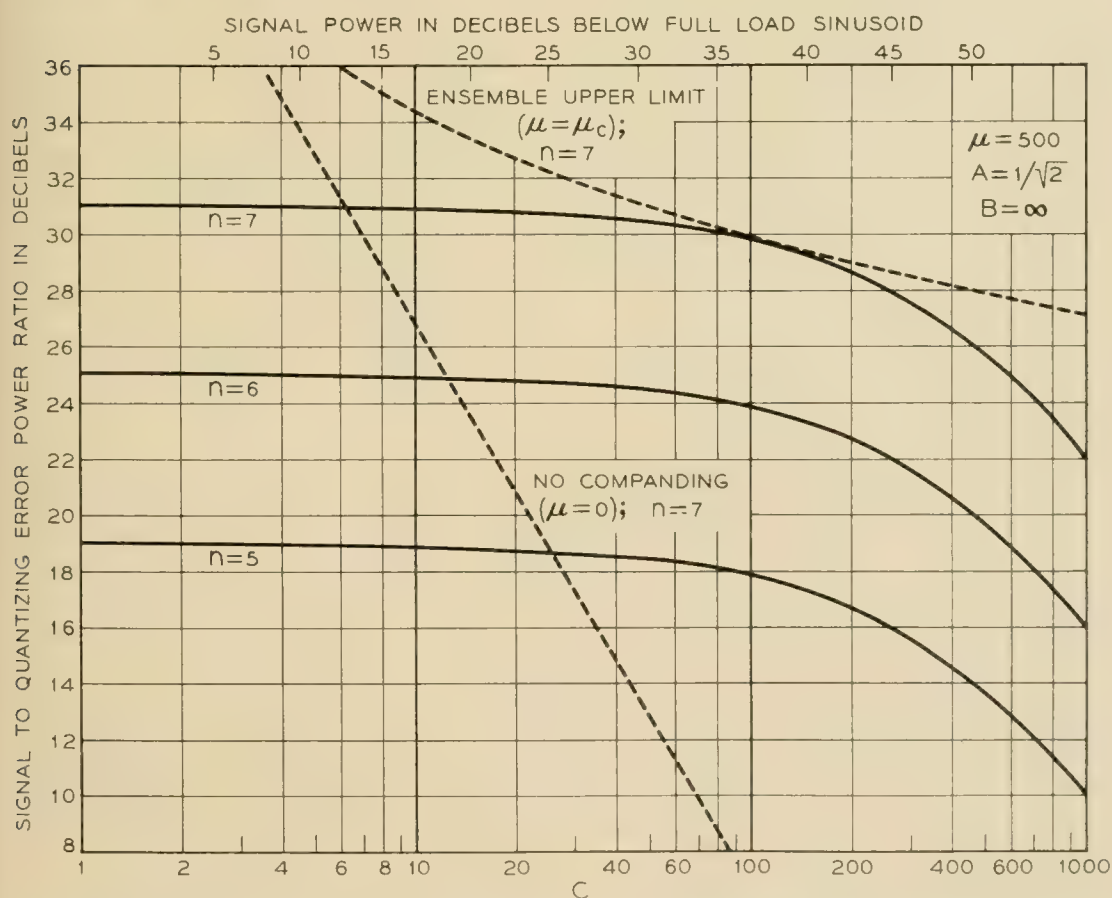


Fig. 17 — Signal to quantizing error power ratios (in db) as a function of relative signal power for companding corresponding to $\mu = 500$. Symbols have the same significance as in Fig. 15. $B = \infty$ throughout.

The recognition that use of $B = 100$ rather than a value approaching 1,000 may imply a change from six to eight digits per code group (e.g., for $\mu = 1,000$), representing an increase of 33 per cent in the required bandwidth in the transmission medium as well as a significant increase in the complexity of the multiplex terminal equipment, provides the proper perspective for competent appraisal of the cost of improving gate circuitry to the point where B would approach 1,000. These considerations might be of crucial importance in the planning of actual PCM systems.

Finally these results also show that caution is required in attempting to determine an adequate number of digits and/or degree of compression from listening tests employing preliminary experimental equipment. If the conditions of the test do not duplicate exactly the expected behavior of the channel gates to be used in the final system, the transition from the laboratory to practice might lead to an embarrassing disappearance

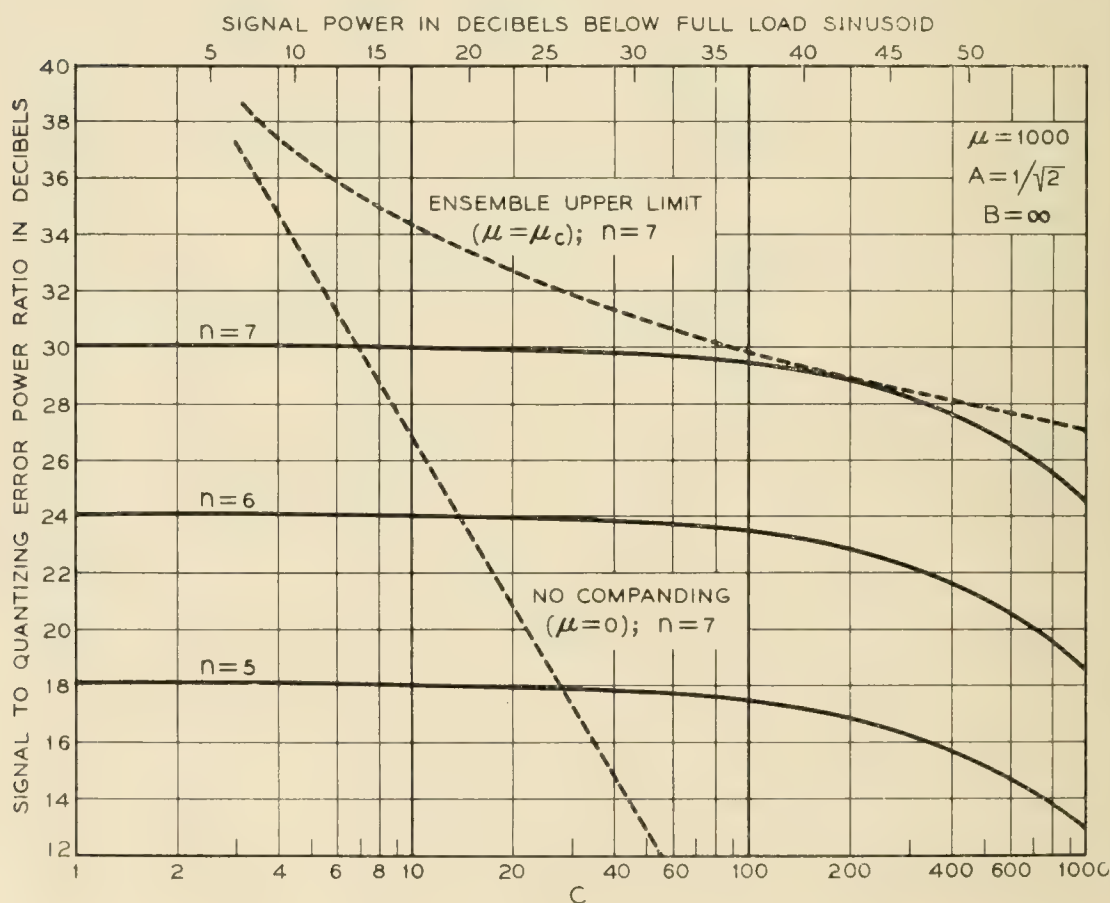


Fig. 18 — Signal to quantizing error power ratios (in db) as a function of relative signal power for companding corresponding to $\mu = 1,000$. Symbols have the same significance as in Fig. 15. $B = \infty$ throughout.

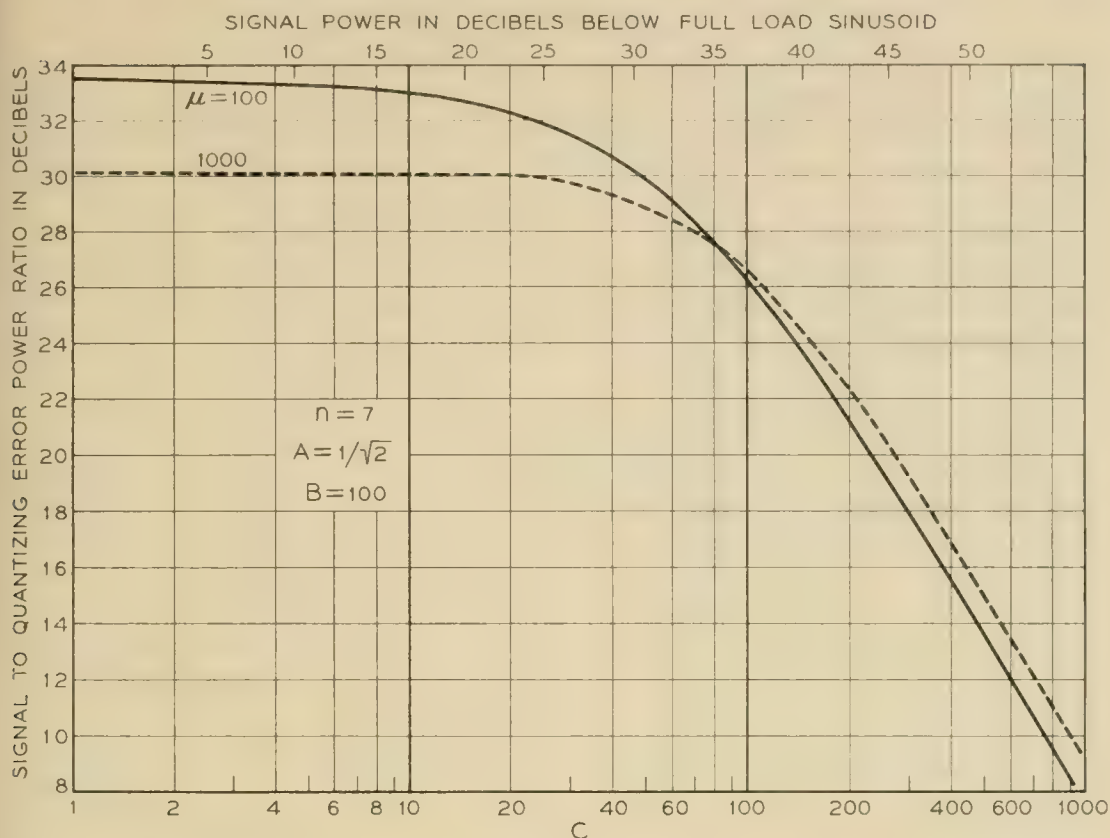


Fig. 19 — Signal to quantizing error power ratios (in db) as a function of relative signal power for companding corresponding to $\mu = 100$ and 1,000 when $n = 7$ digits per code group and a dc component corresponding to $B = 100$ is present in the signal. The influence of the dc component may be judged by comparing these curves with those shown in Figs. 15 and 18 for $n = 7$. Corresponding results for different values of n may be derived by the addition or subtraction of appropriate multiples of 6 db from each ordinate.

of virtually all the anticipated companding improvement for weak signals.

(b) *Background Noise Level.* We have already noted the probable futility of increasing the signal to quantizing error power ratio considerably beyond that value which is subjectively equivalent to the anticipated ratio of signal to background noise from other sources.

If the subjective relation between quantizing error power and noise power were known, the curves in Figs. 15 to 19 could be redrawn for meaningful comparison with ratios of signal to background noise. In the absence of such information, we shall assume as a first approximation, that noise and quantizing error power are directly comparable.*

Suppose that we set an upper limit on the background noise by con-

* The similarity between noise and quantizing error power has often been noted. For example, one may consult references 2, 6 and 12 as well as Appendix I on "Noise in PCM Circuits" in Reference 11. The assumption of direct comparability is also to be found in Reference 4.

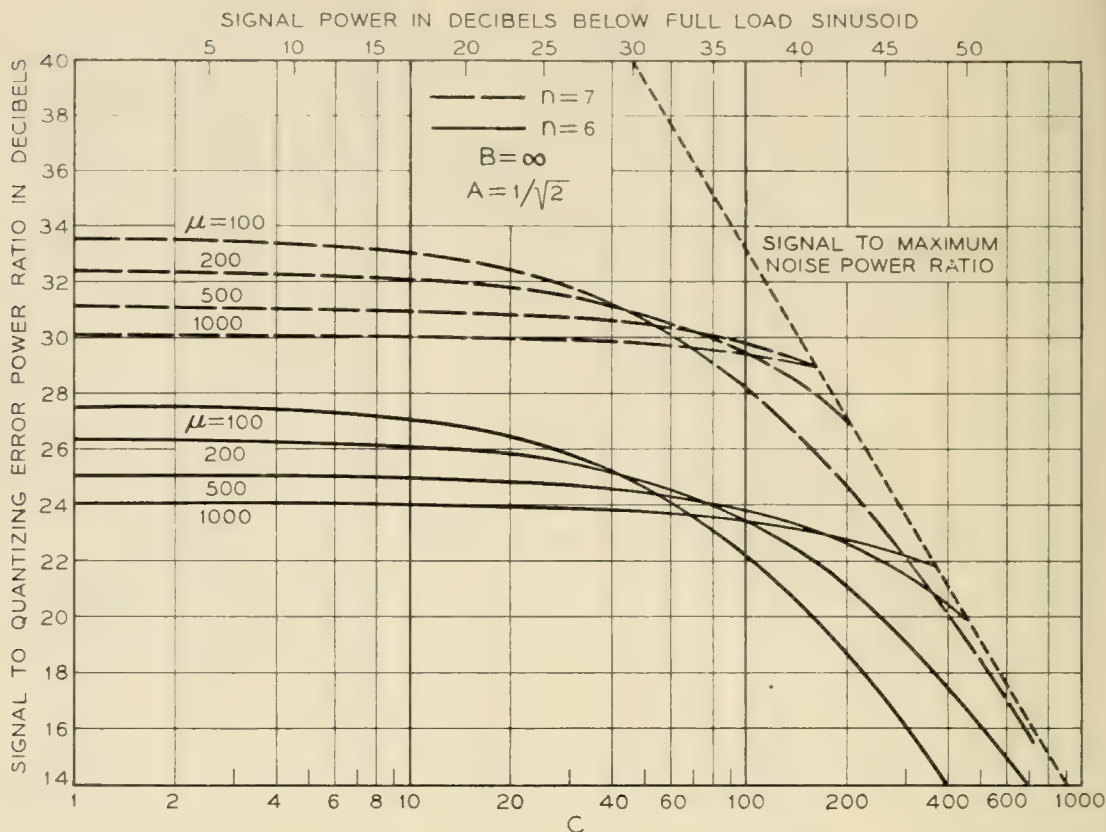


Fig. 20 — Curves illustrating the comparison of signal to quantizing error power ratios with the ratio of signal to background noise. The line representing the signal to maximum noise ratios corresponds to the hypothetical case where the maximum background noise is determined by the requirement that the signal to noise ratio be 20 db for a signal 50 db below full sinusoidal modulation.

sidering a value providing a signal to noise ratio of 20 db for the weakest signals in our hypothetical system. A signal to maximum noise power curve may then be drawn as a function of signal power for this constant value of noise power. Such a graph has been combined, in Fig. 20, with curves such as those which have previously appeared in Figs. 15 to 19. These curves have been terminated at their intersections with the line representing the signal to maximum noise power ratio since we are assuming that little benefit will be derived from a signal to quantizing error power ratio in excess of the signal to maximum noise power ratio.

From Fig. 20 it is apparent that the previous conclusions that six and seven digits are worthy of consideration are unaffected by the stipulation that the signal to quantizing error power ratio should not greatly exceed the signal to maximum noise power ratio. Similarly, the conclusions based on Fig. 19 (for $B = 100$) remain unchanged since the curves therein fall below the maximum noise curve of Fig. 20 for all values of the abscissa.

D. *Possibility of Using Automatic Volume Regulation*

The realization that the quantizing impairment experienced by weak signals in the absence of compression stems from their inability to excite a sufficient number of the quantizing steps which must be provided to accommodate loud signals, leads directly to the suggestion that automatic volume regulation be used to permit all signals to be "loud," i.e., to excite the entire aggregation of quantizing steps. In its simplest form, this would be accomplished by automatic amplification of the long time average speech power in each channel to provide a constant volume input to the common channel equipment.

Study of the present results indicates that if all signals were of constant volume, about 10 to 15 db below full sinusoidal modulation (to provide an adequate peak-clipping margin), satisfactory operation, corresponding to signal to quantizing error power ratios in excess of 20 db, might be achieved *without companding* by using as few as five or six digits per code group. In evaluating this alternative, the advantages of reduction of bandwidth, decreased complexity of quantizing and coding equipment, and elimination of the common channel compandor, must be balanced against the disadvantage of providing separate volume regulators in each channel.

E. *Comparison with Previous Experimental Results*

The literature contains seemingly contradictory statements about whether five,¹⁷ six,¹⁴ or seven^{6, 7} digits per code group are required for satisfactory performance in speech listening tests. Evaluation of these conclusions is frequently hampered by the lack of specification of either the degree of companding employed or the range of speech volumes requiring transmission. Different conclusions may therefore be consistent, inasmuch as the systems may differ significantly in the required volume range, degree of companding, size of the "effective dc component" in the signal, and even in the subjective standards used to judge performance.

Fortunately, the description of an experimental toll quality system by Meacham and Peterson⁶ is sufficiently detailed to permit some comparison. The range of volumes they considered suggests that direct comparison with our hypothetical system is fairly reasonable. Their empirical choice of seven digits, with a compression characteristic virtually indistinguishable from that corresponding to $\mu = 100$ (see Fig. 4) is in excellent agreement with the present conclusions.

Furthermore, the conclusion that five or six digits, without compand-

ing, might be employed in conjunction with volume regulation is completely consistent with Goodall's experimental results.¹⁰

VII. CONCLUSIONS

An effective process for choosing the proper combination of the number of digits per code group and companding characteristic for quantized speech communication systems has been formulated. Under typical conditions, the calculated companding improvement *for the weakest signals* proves to be equivalent to the addition of about 4 to 6 digits per code group, i.e., to an increase in the number of quantizing steps by a factor between $2^4 = 16$ and $2^6 = 64$.

Although a precise application of the results requires a more detailed knowledge of the subjective nature of the quantizing impairment of speech than is presently available, the assumption of reasonably typical system requirements yields conclusions in good agreement with existing experimental evidence.

ACKNOWLEDGMENTS

Frequent references in the text attest to the indebtedness of the author to the writings of Bennett and Panter and Dite. It is also a pleasure to acknowledge stimulating conversations on certain aspects of the problem with J. L. Glaser, D. F. Hoth, B. McMillan, and S. O. Rice.

APPENDIX

THE MINIMIZATION OF QUANTIZING ERROR POWER

In spite of the demonstrated utility of the μ -characteristics, one cannot avoid speculating about the possibility of achieving substantially more companding improvement by using a characteristic which differs from (8). We shall therefore outline a study of the actual minimization of quantizing error power without regard to the relative treatment of various amplitudes in the signal. The results will confirm that a significant reduction of the quantizing error power beyond that attainable with logarithmic companding is self-defeating — for it not only imposes the risk of diminished naturalness, but also implies a compandor too “volume-selective” for the applications envisioned herein.

1. The Variational Problem and Its Formal Solution

Equation (6) may be expressed in the form

$$\sigma = \frac{2V^2}{3N^2} \int_0^V (dv/de)^{-2} P(e) de \quad (\text{A-1})$$

where $P(e)$ has been assumed to be an even function. The function, $v(e)$, which will minimize (A-1), subject to the usual boundary conditions at $e = 0$ and $e = V$, may be obtained by solving the Euler differential equation of the variational problem.²⁸ For (A-1), this takes the form

$$(dv/de) = KP^{1/3} \quad (\text{A-2})$$

where the constant K is given by

$$K = V \bigg/ \int_0^V P^{1/3} de \quad (\text{A-3})$$

Hence the minimum quantizing error is given by

$$\sigma_{\text{MIN}} = 2 \left[\int_0^V P^{1/3} de \right]^3 \bigg/ 3N^2 \quad (\text{A-4})^*$$

2. Representation of Speech by an Exponential Distribution of Amplitudes

We shall assume, as in (25), that the distribution of amplitudes in speech at constant volume¹⁸ may be represented by

$$P(e) = G \exp(-\lambda e) \text{ for } e \geq 0 \quad (\text{A-5})$$

where $P(-e) = P(e)$, $G = \lambda/2$, and $\lambda^2 = 2/\overline{e^2}$. With this choice of $P(e)$, the solution of (A-2)[†] is

$$(v/V) = \frac{1 - \exp [(-\sqrt{2}C/3)(e/V)]}{1 - \exp(-\sqrt{2}C/3)} \quad (\text{A-6})$$

Thus, for any given relative volume (i.e., for each value of $C = V/(\overline{e^2})^{1/2}$), (A-6) specifies the compression characteristic required to minimize the quantizing error power.

We are therefore led to study the properties of the family of characteristics of the form

$$(v/V) = \frac{1 - \exp(-me/V)}{1 - \exp(-m)} \quad \text{for} \quad 0 \leq e \leq V \quad (\text{A-7})$$

* An alternate derivation of (A-2) and (A-4), has been given by Panter and Dite,⁵ who also acknowledge a prior and different deduction by P. R. Aigrain. Upon reading a preliminary version of the present manuscript, B. McMillan called my attention to S. P. Lloyd's related, but unpublished work, which proved to contain still another derivation. I am grateful to Dr. Lloyd for access to this material.

† In the vocabulary of analytical dynamics, the direct integrability of the Euler equation may be ascribed to the existence of an "ignorable" or "cyclic" coordinate.²⁹

TABLE I—COMPARISON OF THE “ μ ” AND “ m ” COMPANDOR ENSEMBLES

Property	μ -Ensemble	m -Ensemble
Defining requirement	Approaches uniform precision in the quantization of all pulse samples outside the unavoidable linear region for small samples	Minimizes quantizing error power for a specific volume, if assume an exponential amplitude distribution for that volume
Compression equation, for $0 \leq e \leq V$, where $v(-e) = -v(e)$	$\left(\frac{v}{V}\right) = \frac{\log\left(1 + \frac{\mu e}{V}\right)}{\log(1 + \mu)}$	$\left(\frac{v}{V}\right) = \frac{1 - \exp(-me/V)}{1 - \exp(-m)}$
Sample to step size ratio	$\left(\frac{e}{\Delta e}\right) = N/2(1 + V/\mu e) \log(1 + \mu)$ $\left(\frac{e}{\Delta e}\right) \rightarrow [N/2 \log(1 + \mu)](\mu e/V)$ for $(e/V) \ll \mu^{-1}$ $\left(\frac{e}{\Delta e}\right) \rightarrow N/2 \log(1 + \mu)$ for $(e/V) \gg \mu^{-1}$	$\left(\frac{e}{\Delta e}\right) = (N/2M)(me/V) \exp(-me/V)$ where $M = 1 - \exp(-m)$ $\left(\frac{e}{\Delta e}\right) \rightarrow (N/2M)(me/V)$ for $(e/V) \ll m^{-1}$ $\left(\frac{e}{\Delta e}\right) = \text{MAX at } (e/V) = m^{-1}$
Ratio of largest to smallest step size = ratio of compression characteristic slopes for zero and overload inputs	$\frac{(\Delta e)_{e=V}}{(\Delta e)_{e=0}} = \frac{(dv/de)_{e=0}}{(dv/de)_{e=V}} = \mu + 1$	$\frac{(\Delta e)_{e=V}}{(\Delta e)_{e=0}} = \frac{(dv/de)_{e=0}}{(dv/de)_{e=V}} = \exp(m)$

Saturated companding improvement, corresponding to linear compression of weakest signals	$C \gg \mu: 20 \log_{10} \left[\frac{\mu}{\log (1 + \mu)} \right] db$	$C \gg m: 20 \log_{10} \left[\frac{m}{1 - \exp (-m)} \right] db$
Relation between relative signal strength and critical compression parameter which provides greatest reduction of quantizing error for a given volume	$C^2 S + \mu_c A C (S - 1) - \mu_c^2 = 0$ where $C = V / \sqrt{e^2} = \text{relative signal strength}$ $S = \left(\frac{1 + \mu_c}{\mu_c} \right) \log (1 + \mu_c) - 1$ $A = e / \sqrt{e^2} = \text{indep. of volume}$	$C = 3m_c / \sqrt{2}$
$D^2 = \frac{\text{Error Power}}{\text{Signal Power}}$ when $\bar{e} = 0$	$D^2 = \frac{[\log (1 + \mu)]^2}{3N^2} \left[1 + \frac{C^2}{\mu^2} + \frac{2AC}{\mu} \right]$	$D^2 = \frac{\sqrt{2} C^3 M^2}{3N^2 m^2} \left[\frac{\exp (2m - \sqrt{2} C) - 1}{2m - \sqrt{2} C} \right]$

with $v(-e) = -v(e)$ as usual. The “ m -characteristics” specified by (A-7) are to be compared with the “ μ -characteristics” specified by (8).

From the derivation of (A-6) it is known that optimum companding will be produced when m is given by the critical value,

$$m_c = \sqrt{2C}/3 \quad (\text{A-8})$$

This is the analogue of (36) defining μ_c for the μ -ensemble.

3. Properties of the “ m -Ensemble”

We shall now interpret the properties of the m -ensemble of companders, for which the ensemble improvement limit ($m = m_c$) actually

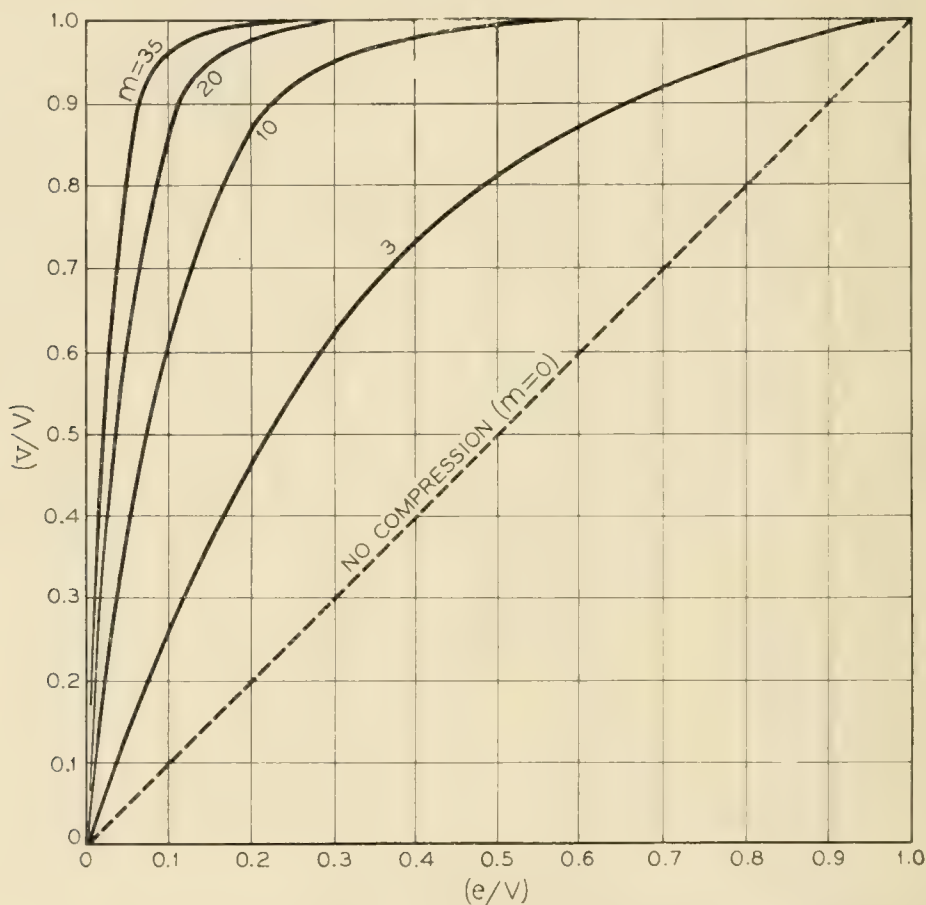


Fig. 21 — Typical m -ensemble compression characteristics determined by equation (A-7). Note the strong emphasis on weak signal amplitudes. These curves may be compared with those for the μ -ensemble in Fig. 3.

minimizes the total quantizing error power, when the probability density is specified by (A-5). Table I summarizes the important properties which may be derived by replacing (8) by (A-7) in the previous detailed analysis of the μ -ensemble.

(a) *Compression Characteristics*

Compression characteristics, corresponding to various values of m are displayed in Figs. 21 and 22 for direct comparison with the curves in Figs. 3 and 4. The m -characteristics assign very little weight to the larger signal amplitudes in view of the infrequent occurrence of the latter.

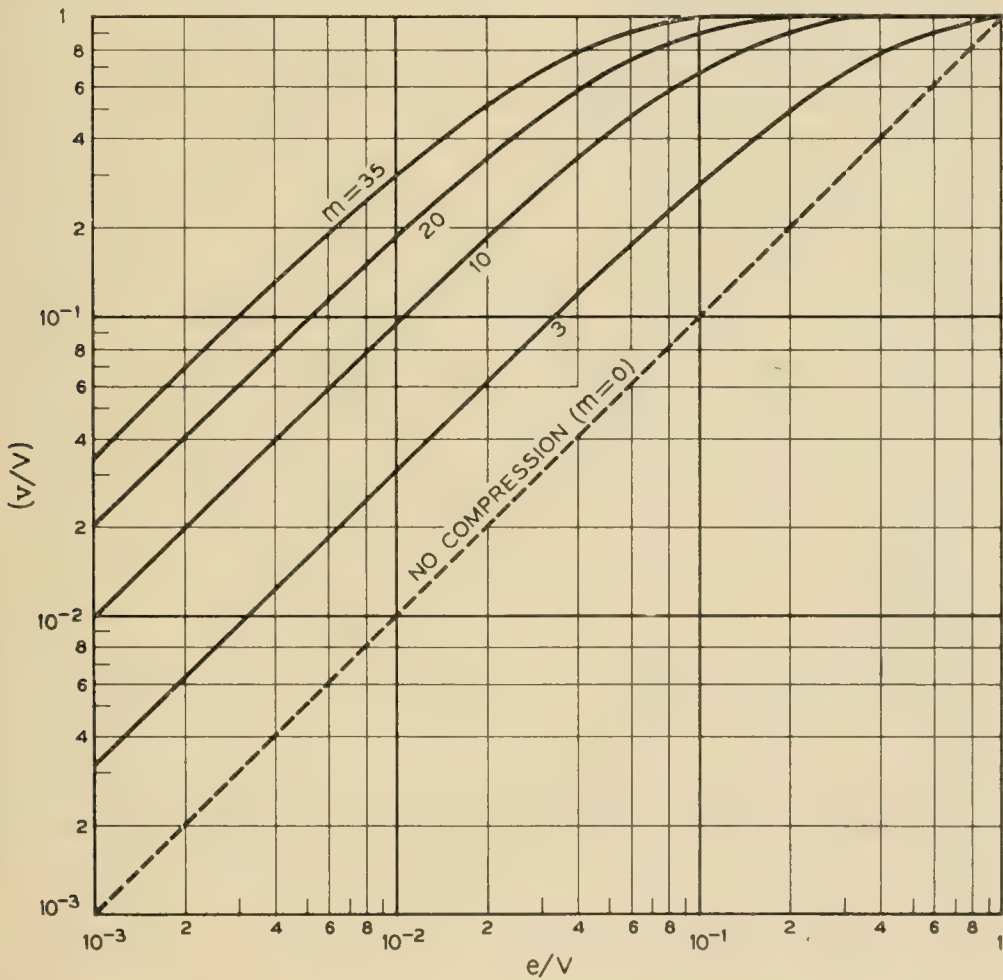


Fig. 22 — Logarithmic replot of compression curves of the type shown in Fig. 21 to indicate detailed behavior for weak samples. These may be compared with the μ -ensemble curves in Fig. 4.

(b) *Sample to Step Size Ratio*

Fig. 23, where the sample to step size ratio ($e/\Delta e$) is plotted in the same manner as in Fig. 5, reveals the relative quantizing accuracy accorded various pulse amplitudes.

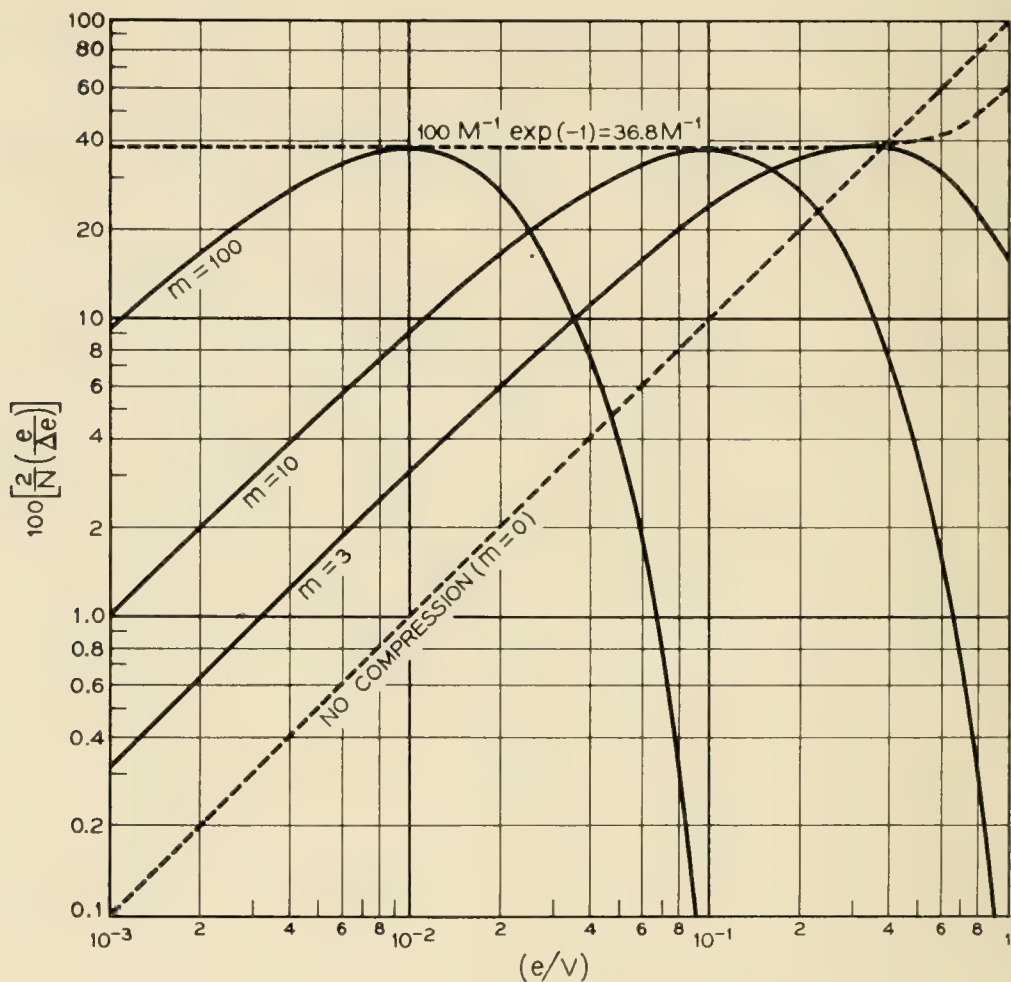


Fig. 23 — Pulse sample to step size ratios as a function of relative sample amplitude, for various companders in the m -ensemble. The maxima exhibited by these curves occur at $e/V = m^{-1}$; $M = 1 - \exp(-m)$. Compare with Fig. 5.

(c) *Saturated Improvement of Weak Signals*

For signals whose largest samples are confined to the region $(e/V) \ll m^{-1}$, compression is linear, with a saturation improvement noted in Table I and plotted in Fig. 24 for comparison with Fig. 9.

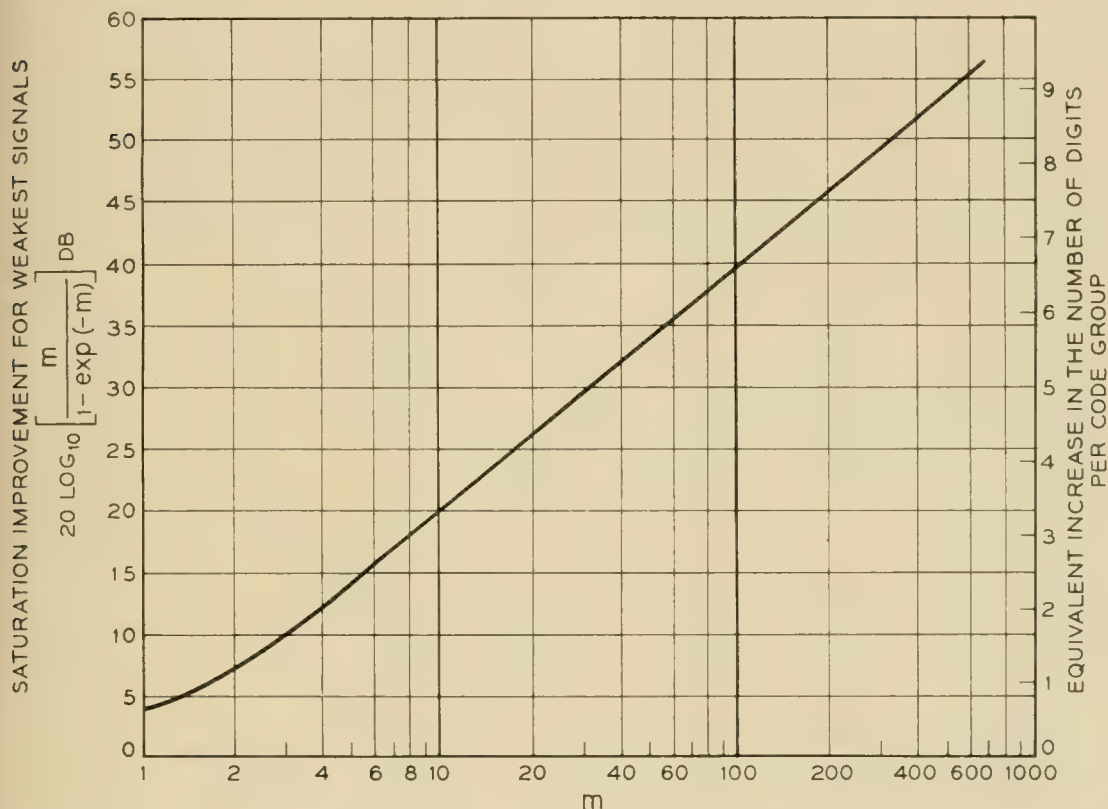


Fig. 24 — Saturated companding improvement for the weakest signals as a function of the degree of “ m -type” compression. Given a value of m , the corresponding ordinate represents the reduction of quantizing error power (in db) which results from companding of signals so weak that signal peaks satisfy the relation $(e/V) \ll m^{-1}$. Thus, weaker and weaker signals are required to exploit the added improvement which follows from an increase in m . This curve may be compared with that for the μ -ensemble in Fig. 9.

(d) *Variation of Companding Improvement with Volume*

Companding improvement curves are shown in Fig. 25 for representative members of the m -ensemble. Each curve is tangent to the ensemble upper limit at the volume for which $m = m_c$. In view of its deduction as the solution of the variational problem, this upper limit actually represents the absolute maximum value of companding improvement for the present choice of $P(e)$.

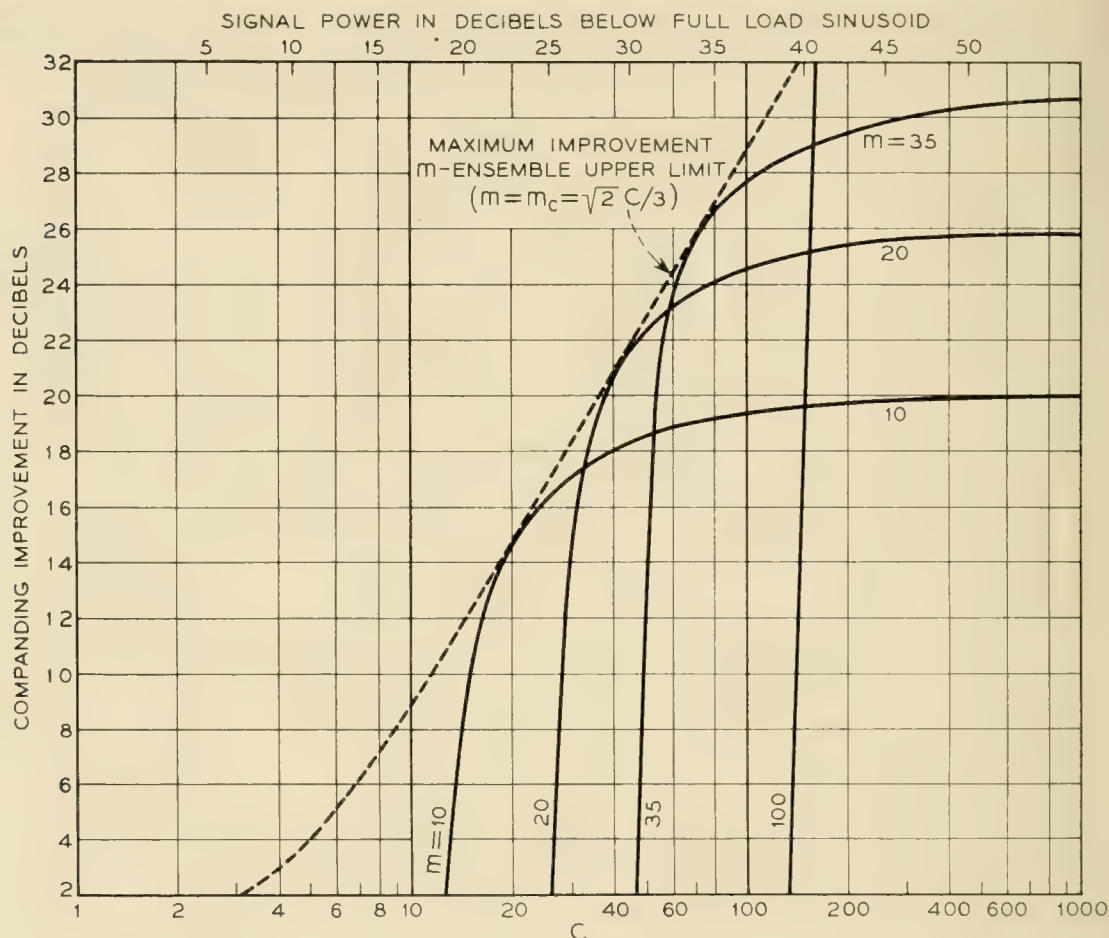


Fig. 25 — Companding improvement curves for representative members of the m -ensemble. These curves are to be compared with those for the μ -ensemble in Fig. 8. Note the important difference between the two ensembles for strong signals (small values of C).

(e) Signal to Quantizing Error Power Ratios

The curves in Fig. 26 are drawn for the representative case of $N = 2^7 = 128$ quantizing steps (7 digit PCM). The corresponding ensemble limit is constant, as might be expected from (A-4), except for strong signals where the effects of peak clipping become noticeable.

In the region where this ensemble limit is constant, departures from the improvement limit resulting from the use of a single value of m for all volumes may be read directly from the ordinates shown at the right in Fig. 26. In comparing these departures from maximum improvement with the analogous μ -ensemble curves in Fig. 14, it must always be recalled that, in view of its role in the solution of the variational problem, the m -ensemble limit represents the actual minimum quantizing error power consistent with the probability density specified by (A-5).

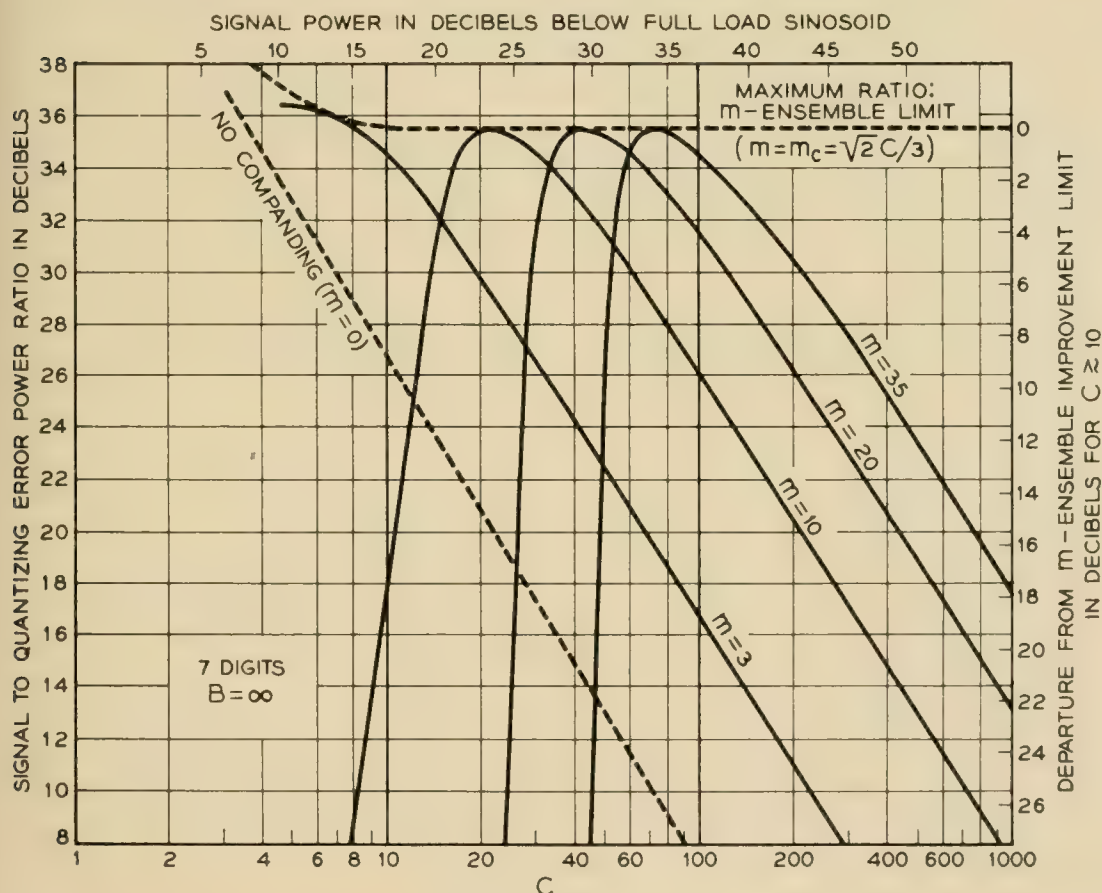


Fig. 26 — Signal to quantizing error power ratios as a function of relative signal power for 7 digits and various m -compandors. The curves may be compared with those for 7 digits in Figs. 15 to 18. The auxiliary ordinates at the right of the present figure apply for $C \gtrsim 10$, where the m -ensemble limit is effectively constant; departures from this limit, resulting from the use of a single value of m for all volumes, may be read directly from this scale, for comparison with Fig. 14. The latter comparison illustrates the narrow volume limitation of the members of the m -ensemble.

(f) Illustrative Application

Consider the possibility of choosing a member of the m -ensemble for application to the hypothetical PCM system already discussed in connection with the μ -ensemble. It will be recalled (see Figs. 15–18) that we were able to choose degrees of logarithmic compression which would yield signal to quantizing error power ratios in excess of about 20 db for all volumes ($4.5 \lesssim C \lesssim 450$) by using as few as six or seven (depending on the choice of μ) digits per code group. In contrast, Fig. 26 reveals that no value of m will meet this requirement since the curves fall so rapidly on either side of the sharp maxima. In short, the members of the m -ensemble are each too specialized for successful application to such a broad volume range.

Further detailed comparison between the numerical results for the two ensembles seems inappropriate, since it is not at all clear that the inequitable treatment of the various samples in a given signal by members of the m -ensemble (see Fig. 23) permits an adequate description of signal quality solely in terms of quantizing error power. Under these circumstances, subjective effects beyond the scope of the present analysis might assume a dominant role.

REFERENCES

1. H. S. Black, *Modulation Theory*, Van Nostrand, N. Y., 1953.
2. W. R. Bennett, Spectra of Quantized Signals, *B.S.T.J.*, **27**, pp. 446–472, July, 1948.
3. Oliver, Pierce and Shannon, The Philosophy of PCM, *Proc. I.R.E.*, **36**, pp. 1324–1331, Nov., 1948.
4. Clavier, Panter and Dite, Signal-to-Noise-Ratio Improvement in a PCM System, *Proc. I.R.E.*, **37**, pp. 355–359, April, 1949.
5. P. F. Panter and W. Dite, Quantization Distortion in Pulse-Count Modulation with Nonuniform Spacing of Levels, *Proc. I.R.E.*, **39**, pp. 44–48, Jan., 1951.
6. L. A. Meacham and E. Peterson, An Experimental Multichannel Pulse Code Modulation System of Toll Quality, *B.S.T.J.*, **27**, pp. 1–43, Jan., 1948.
7. H. S. Black and J. O. Edson, Pulse Code Modulation, *Trans. A.I.E.E.*, **66**, pp. 895–899, 1947.
8. Clavier, Panter and Grieg, Distortion in a Pulse Count Modulation System, *Elec. Eng.*, **66**, pp. 1110–1122, 1947.
9. J. P. Schouten and H. W. F. Van' T. Groenewout, Analysis of Distortion in Pulse Code Modulation Systems, *Applied Scientific Research*, **2B**, pp. 277–290, 1952.
10. W. M. Goodall, Telephony by Pulse Code Modulation, *B.S.T.J.*, **26**, pp. 395–409, July, 1947.
11. C. B. Feldman and W. R. Bennett, Bandwidth and Transmission Performance, *B.S.T.J.*, **28**, pp. 490–595, July, 1949.
12. W. R. Bennett, Sources and Properties of Electrical Noise, *Elec. Eng.*, **73**, pp. 1001–1008, Nov., 1954.
13. C. Villars, Etude sur la Modulation par Impulsions Codées, *Bulletin Technique PTT*, pp. 449–472, 1954.
14. J. Boisivieux, Le Multiplex à 16 Voies à Modulation Codée de la C.F.T.H., *L'Onde Electrique*, **34**, pp. 363–371, Apr., 1954.

15. E. Kettel, Der Störabstand bei der Nachrichtenübertragung durch Codemodulation, *Archiv der Elektrischen Übertragung*, **3**, pp. 161-164, Jan., 1949.
16. H. Holzwarth, Pulsecodemodulation und ihre Verzerrungen bei logarithmischer Amplitudenquantelung, *Archiv der Elektrischen Übertragung*, **3**, pp. 277-285, Jan., 1949.
17. Herreng, Blondé and Dureau, Système de Transmission Téléphonique Multiplex à Modulation par Impulsions Codées, *Cables & Transmission*, **9**, pp. 144-160, April, 1955.
18. W. B. Davenport, Jr., An Experimental Study of Speech-Wave Probability Distributions, *J. Acous. Soc. Amer.*, **24**, pp. 390-399, July, 1952.
19. S. B. Wright, Amplitude Range Control, *B.S.T.J.*, **17**, pp. 520-538, Oct., 1938.
20. Carter, Dickieson and Mitchell, Application of Companders to Telephone Circuits, *Trans. A.I.E.E.*, **65**, pp. 1079-1086, Dec., 1946.
21. J. C. R. Licklider and I. Pollack, Effects of Differentiation, Integration and Infinite Peak Clipping upon the Intelligibility of Speech, *J. Acous. Soc. Amer.*, **20**, pp. 42-51, Jan., 1948.
22. D. W. Martin, Uniform Speech-Peak Clipping in a Uniform Signal to Noise Spectrum Ratio, *J. Acous. Soc. Amer.*, **22**, pp. 614-621, Sept., 1950.
23. J. C. R. Licklider, The Intelligibility of Amplitude-Dichotomized, Time-Quantized Speech Waves, *J. Acous. Soc. Amer.*, **22**, pp. 820-823, Nov., 1950.
24. H. Cramér, *Mathematical Methods of Statistics*, Princeton Univ. Press, Princeton, N. J., 1946, see pp. 359-363.
25. T. C. Fry, *Probability and its Engineering Uses*, Van Nostrand, N. Y., 1928, see pp. 310-312.
26. A. C. Aitken, *Statistical Mathematics*, Interscience Pub. Inc., N. Y., 3d Ed., 1944, see pp. 44-47.
27. W. F. Sheppard, On the Calculation of the Most Probable Values of Frequency-Constants for Data Arranged According to Equidistant Divisions of a Scale, *Proc. London Math. Soc.*, **29**, pp. 353-380, 1898.
28. R. Courant and D. Hilbert, *Methods of Mathematical Physics*, Vol. 1, Interscience Pub. Inc., N. Y., English Ed., 1953, see Chapt. IV, especially pp. 184-187, and p. 206.
29. E. T. Whittaker, *A Treatise on the Analytical Dynamics of Particles and Rigid Bodies*, Cambridge Univ. Press, 4th Ed., 1937, see p. 54.

W. D. Bulloch Appointed Editor of B.S.T.J.

W. D. Bulloch, formerly Editor of the Bell Laboratories Record, has been appointed Editor of the Bell System Technical Journal.

Mr. Bulloch received a bachelor's degree from Dartmouth College and a Master of Science degree in Physics from the University of North Carolina. He taught physics, mathematics and astronomy in the latter institution for several years before joining the staff of Bell Telephone Laboratories.

An Electrically Operated Hydraulic Control Valve

By J. W. SCHAEFER

(Manuscript received August 3, 1956)

The electrohydraulic transducer used in the servos that drive the control surfaces of the NIKE missile is described and its operating characteristics are discussed. Special attention is directed to the secondary dynamic forces that exist in a high-gain device of this type and to the resulting tendency to oscillate. The application of the valve to a servo system is discussed briefly.

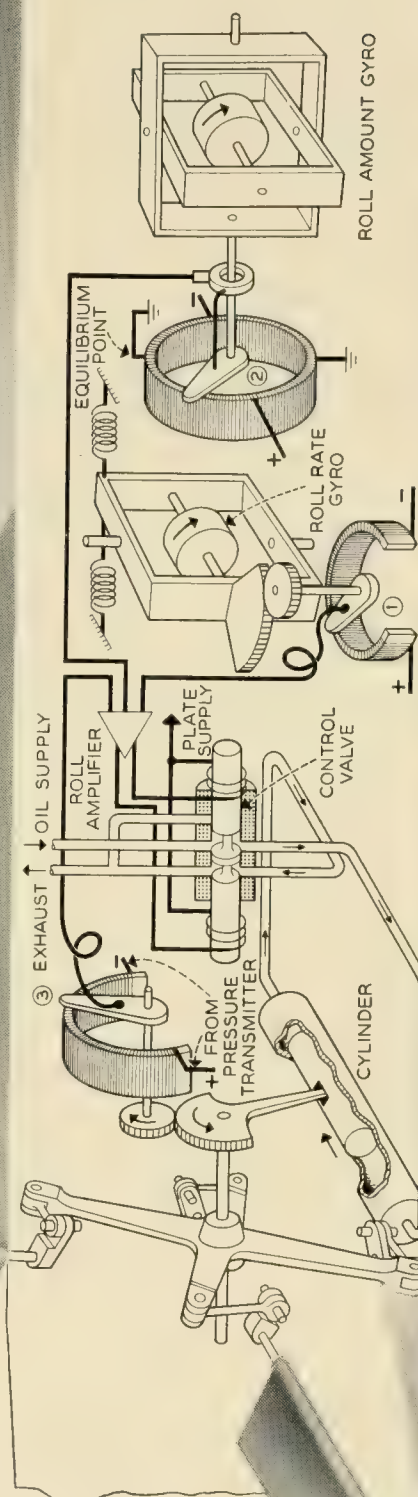
INTRODUCTION

Early in the study of the NIKE guided missile project, it became apparent that the requirements for the fin actuators could not be fulfilled by the servo-mechanisms available at that time (1945). All existing types failed to meet the combined requirements of small size, light weight, high torque, and rapid response. Further investigation showed that the development of a hydraulic servo employing an electrohydraulic transducer appeared to provide a promising solution. A control system of this type, therefore, has been developed for the NIKE missile.

The design of the transducer, or control valve, was one of the principal problems in the development of the missile control systems and is the subject of this article. The specific design of the valve that will be discussed here is known as "Model J-7", and represents the state of the development in 1950. Valves of this type, with varying degrees of modification, are used in missiles of several other projects.

APPLICATION

Fig. 1 is a simplified schematic of the roll positioning system in the NIKE missile. It is the simplest of the three applications of the valve in the missile, but will serve to illustrate the situation in which the valve operates. The purpose of the roll servo is to keep the missile in a predictable roll orientation.



The roll system's reference is an "Amount Gyro," which is a free-free gyro oriented on the ground prior to missile launch. The brush of a four-tap potentiometer (Item 2 in Fig. 1) is connected to the outer gimbal and provides a dc signal whose sign and magnitude indicate the roll position with respect to the stable equilibrium point. This signal is the principal input to the servo amplifier that drives the valve.

A roll-position error exists in the situation illustrated in Fig. 1. The valve is driven in the direction to cause the oil flow to rotate the ailerons, which in turn will roll the missile toward the null position. As the missile rolls, the winding of the roll-amount potentiometer rotates with it. The brush stays fixed in space with the gyro gimbal.

The aerodynamic coupling between the aileron position and the missile's roll position is a complex and variable term in the feedback loop of the servo. The nature of the aerodynamic coupling is such that an otherwise simple servo problem becomes considerably more complicated. During a normal flight the aerodynamic stiffness, and hence the gain in the feedback loop, varies over a 50:1 range. A first order correction for this change is accomplished by a variable gain local loop around the valve, cylinder and amplifier. A potentiometer (Item 3 in Fig. 1) is geared to the fin in such a way that a dc signal is produced, which is proportional to fin position. The gain of this local loop is varied by supplying the potentiometer with voltages that are directly proportional to the measured aerodynamic stiffness. In this way the amount of the deflection of the aileron is made inversely proportional to the aerodynamic stiffness. This effect results in an approximately constant torque about the roll axis of the missile for a given signal. The local loop around the fin position also reduces the effect of any non-linear characteristics of the valve.

A third input to the servo amplifier is provided by a potentiometer that is driven by a spring-restrained gyroscope mounted so that the sensitive axis is aligned with the missile's roll axis. The dc signal produced is proportional to the roll rate. This signal provides some anticipation to the roll position loop. It performs the function of a tachometer in a conventional servo. It also insures that the roll rate is limited to a value that can be handled by the position loop. If very high roll rates were allowed to exist, the roll amount gyro would produce a signal changing in sign at such a rate that the ailerons would be unable to keep up or reduce the rate.

The roll servo insures that the missile's orientation is aligned with the free-free gyro. This enables the yaw and pitch servos to steer about their assigned axes in a consistent manner. The two steering servos are

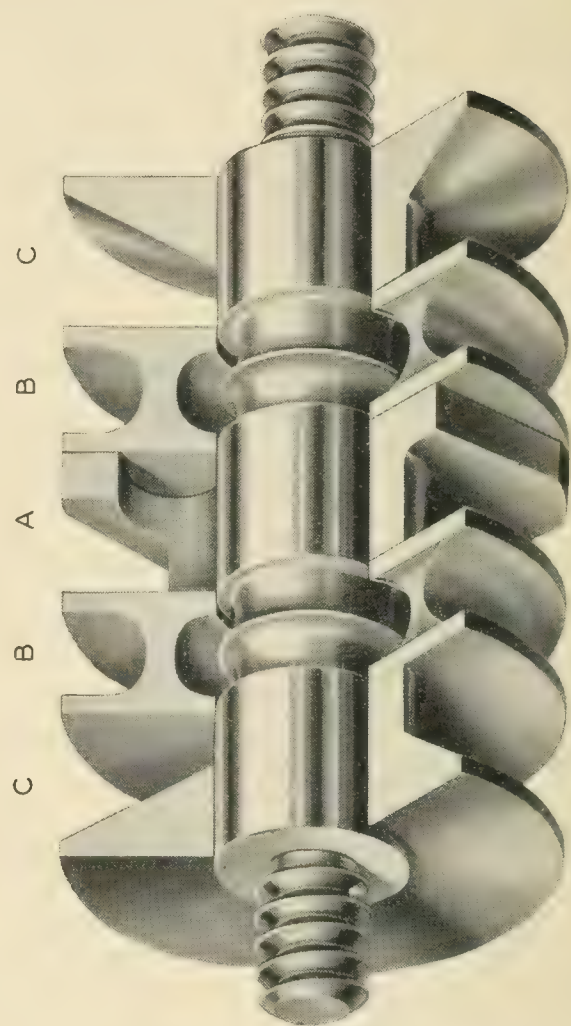
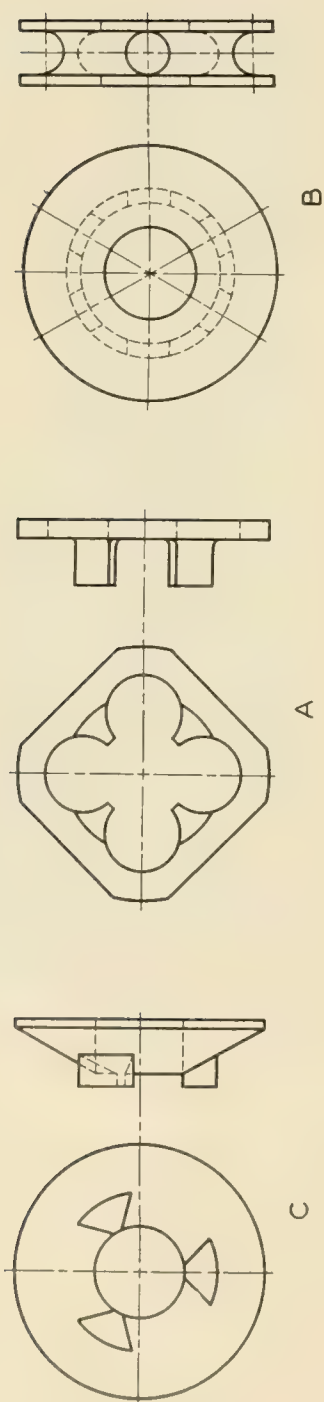
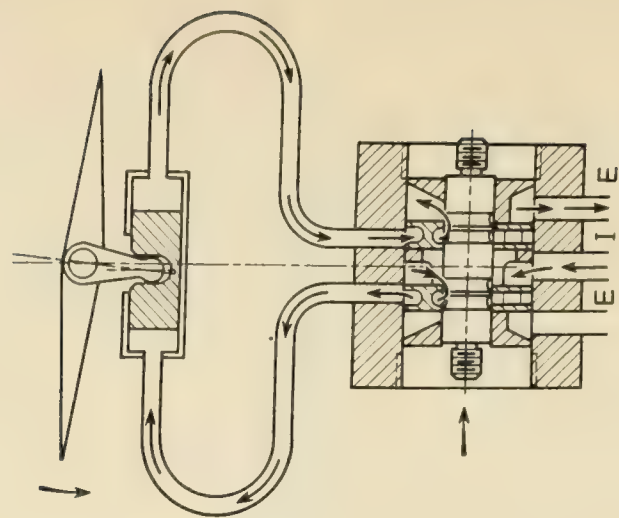


Fig. 2 — Porting arrangement, J-7 solenoid valve.

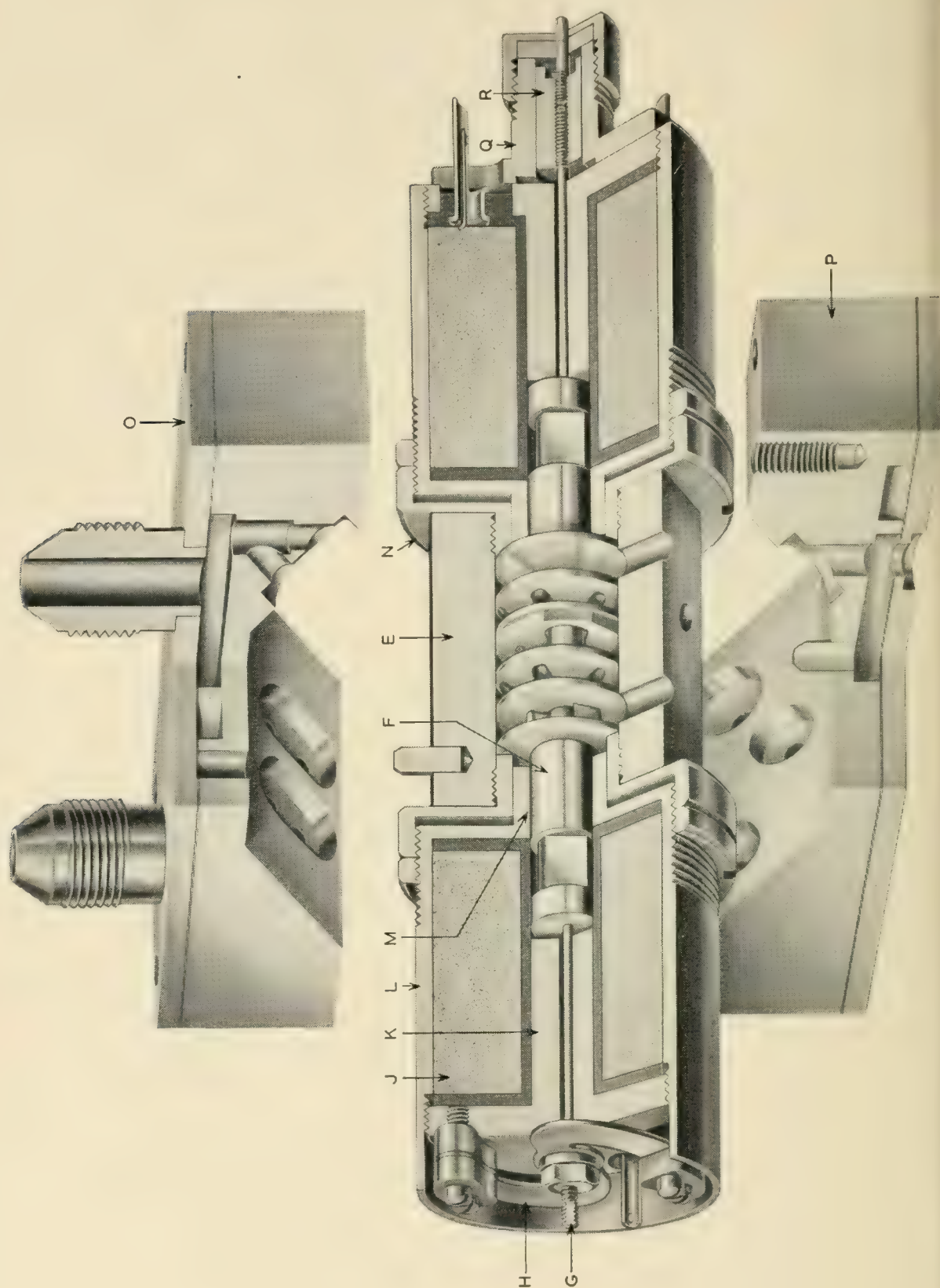
identical to each other and somewhat similar to the roll system described above. Each of the three systems employs identical valves.

GENERAL DESCRIPTION

Basically, the J-7 valve is a conventional four-way type. Fig. 2 illustrates the porting arrangement. The parts in sections A, B and C are inserts that are shrunk-fit into the valve body. The plunger accurately fits the holes in the inserts, so that oil cannot flow between the plunger and inserts except where the diameter of the plunger is reduced. The annular space around the outside of the center insert is connected to the high pressure oil supply. The radial passages in this insert (part A) carry the oil to its internal cusps. With the plunger centrally located its center land completely covers the port formed by the cusps and no oil flows. With the plunger moved to the right the oil is carried to the cylinder and back to the exhaust in the manner illustrated by the small sketch in the upper right corner. If the motion of the plunger is to the left, a similar performance occurs but the piston and fin are driven in the opposite direction. To illustrate the construction of the inserts, detail sketches are also shown at the top of Fig. 2.

The inserts and the plunger are made of hardened steel. Fig. 3 shows their location in the complete valve. The thickness of the inserts, hence the longitudinal location of the ports, is held to an extremely close tolerance by lapping their parallel faces. Their outside diameter is accurately ground so that a tight seal will occur between the various passages when they are shrunk fit into the internal bore of the body. After assembly, the internal bore formed by the holes in the various inserts is lapped to a straight and accurate cylindrical shape. This process is controlled to provide a diametral clearance of 0.0002 inch on an interchangeable basis. The plunger must slide freely in the bore in spite of the small clearances involved. The longitudinal location of the lands on the plunger must be controlled to a high degree for reasons that will become apparent.

Those parts shown in Fig. 2 are not sectioned in Fig. 3. The valve proper is clamped between two manifolds (O and P); these are moved apart in the picture to better illustrate the internal construction. The brazed laminated manifolds provide the mounting means for the valve and also serve to connect the multiple outlets of the valve body to standard hydraulic fittings for external connections. The manifolds are designed to adapt the valve to a specific application. In this way, different plumbing arrangements can be utilized without changes in the valve proper. The manifold, O, has fittings to connect to the cylinder,



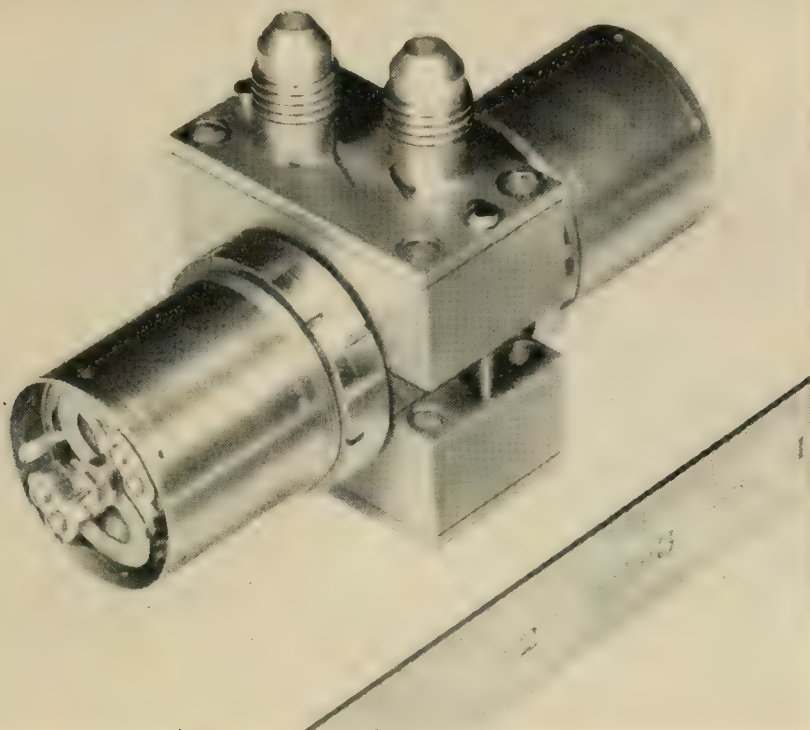


Fig. 4 — J-7 solenoid valve with manifolds.

and the lower manifold, P, serves to connect the pressure and exhaust lines to ports in the structure to which the valve is mounted. The joints between the passages in the manifolds and those in the body are sealed by rubber "O" rings that are inserted in the recesses about the holes on the inner faces of the manifolds. These recesses can be seen in Fig. 3, but the "O" rings are not illustrated. Similarly, "O" rings are used to seal the joints between the manifold, P, and the flat surface to which it mounts.

A pole piece, *F*, is screwed to each end of the plunger by means of the threads visible in Fig. 2. These parts move as an assembly and form

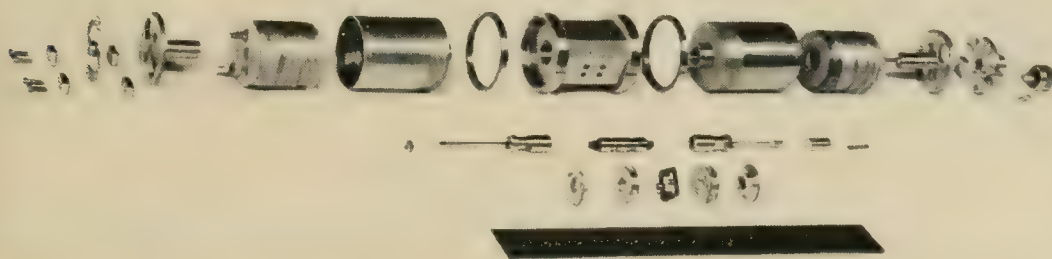


Fig. 5 — Exploded view of J-7 solenoid valve.

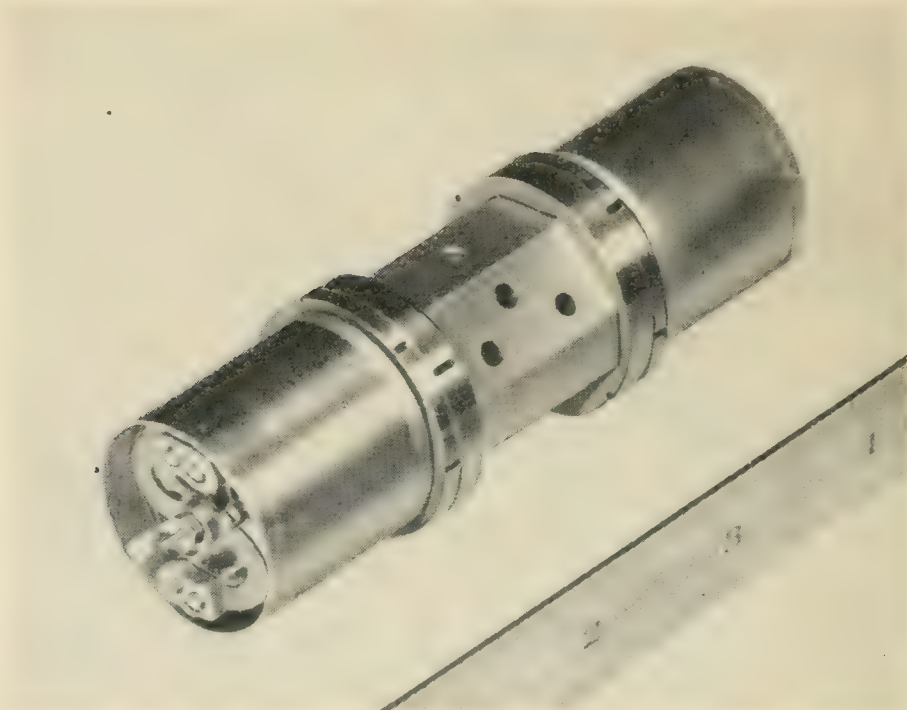


Fig. 6 — J-7 solenoid valve.

the armature of the valve. The push rod, G, attached to the armature, is connected to an S shaped spring, H, which tends to keep the movable assembly centered, or the valve closed.

When a current passes through the coil, J, magnetic flux passes through the fixed pole piece, K, through the coil housing, L, then across the small annular air gap, M, to the moving pole piece, F, and across the air gap between the pole faces near the center of the coil. Flux in the latter gap causes a force on the armature which tends to pull it and the valve

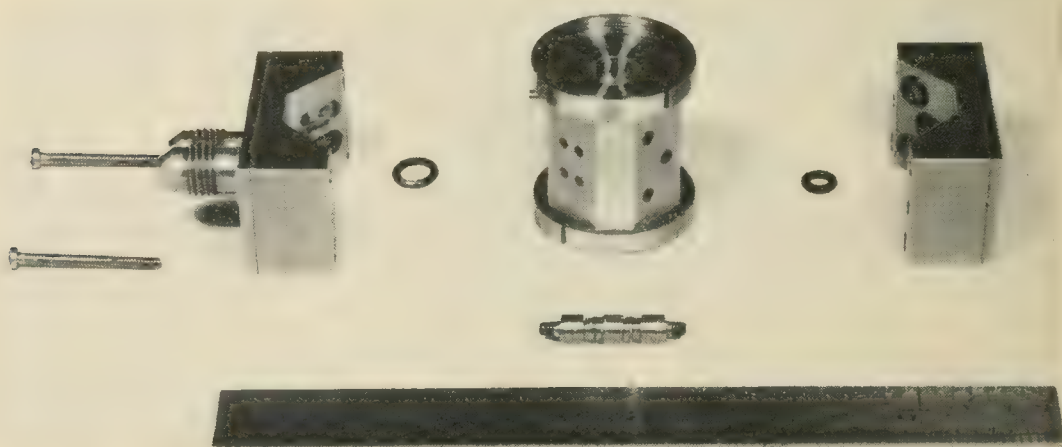


Fig. 7 — Hydraulic parts of the J-7 valve.

plunger toward that coil. If an equal current is flowing in the other coil the forces are balanced and the valve remains centered and closed. If the currents in the two coils are not equal, the valve plunger is moved until the differential magnetic force is balanced by the force in the spring. In this manner the amount of oil flow can be regulated by varying the difference of the currents in the two coils.

The coil housing is attached to the valve body by means of a non-magnetic stainless steel adapter, N. The adapter isolates the steel plunger and inserts from the magnetic flux. Because of the close fit between these parts the presence of flux would cause sticking. A push rod attached to the armature drives an aluminum piston, R, in a cylinder, Q. The small radial clearance between these parts is filled with a viscous fluid so that a damping force is produced that is proportional to armature velocity.

Fig. 4 shows the complete assembly with manifolds attached, while Fig. 5 gives an exploded view of the valve proper with all details in their correct relative positions. Other views of the valve parts are shown in Figs. 6 and 7. Fig. 8 is a view looking into the bore of the hydraulic assembly. This is the hole that is normally occupied by the plunger. The ports formed by the internal shapes of the inserts are clearly visible.

CHARACTERISTICS OF THE ACTUATING MECHANISM

The J-7 valve is designed to be driven by a push-pull dc amplifier. When the amplifier has no input signal, the output current in each side is 10 ma. When a signal is applied, the current in one side is increased and that in the other is decreased; at maximum signal, the current in one coil reaches 19 ma and zero in the other. The dissipation in each coil for quiescent current is about 0.4 watt, and the full signal power is 1.5 watts.

The requirements that the frequency response must extend to dc and that the control power consumption be held to a minimum suggest the use of a dc push-pull output stage. The quiescent dc plate current is used as the magnetizing current for the solenoids instead of providing the field by a permanent magnet or separate coil. In spite of the small output current available from the amplifier, relatively large forces and a high resonant frequency are realized. This is accomplished by an efficient magnetic circuit and low armature mass. The opposing solenoid configuration described makes these features possible.

The magnetic circuit used in the J-7 valve has very low reluctance for a sliding armature type actuator. This low reluctance is accomplished by providing ample thickness in the iron parts and employing

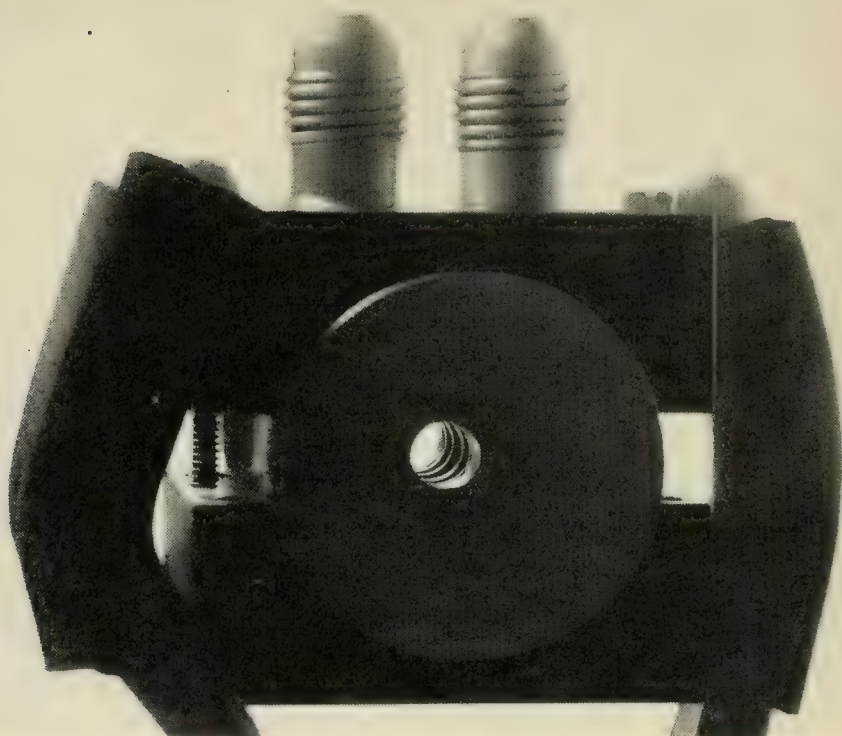


Fig. 8 — Internal view of the J-7 valve body.

very short gaps of considerable cross-sectional areas. It has been shown elsewhere^{1, 2} that if the distribution of reluctance of the magnetic structure is symmetrical along its longitudinal axis, minimum leakage flux and hence maximum force would be developed if the working gap were at the exact center of the coil. To a close approximation, the valve solenoids are symmetrical in this manner. The references show that the maximum pull is not sensitive to small changes from the optimum location of the working gap. The gap in the J-7 valve is displaced toward the center of the valve from the location of maximum pull. The slight reduction in magnetic force was accepted as a suitable compromise for the resulting reduction in the mass of the moving pole piece and the corresponding increase in resonant frequency.

The gap between the solenoid pole faces is 0.014 inch when there is no signal and the valve is centered. The two fixed faces have 0.004 inch non-magnetic shims attached so that the maximum motion of the armature is limited to ± 0.010 inch. The shims prevent the armature from sticking against the fixed faces by maintaining appreciable gaps. To further compensate for the inverse square law of magnetic attraction, the moving pole pieces have a reduced section that saturates under high flux or large forces. This neck can be seen plainly in Fig. 3. The saturating

sections limit the flux and tend to reduce the pull for short gaps so that the centering spring can be a simple linear member and still not lose control when the armature is near the fixed pole face. The flat surfaces on the neck permit the use of wrenches for assembly. Placing the necks on the moving pole pieces further reduces the mass of the armature assembly.

To illustrate the saturation action, Fig. 9 shows the pull of one of the magnets plotted against gap length for quiescent and maximum current. The shim line and curves from a solenoid without a saturation neck also

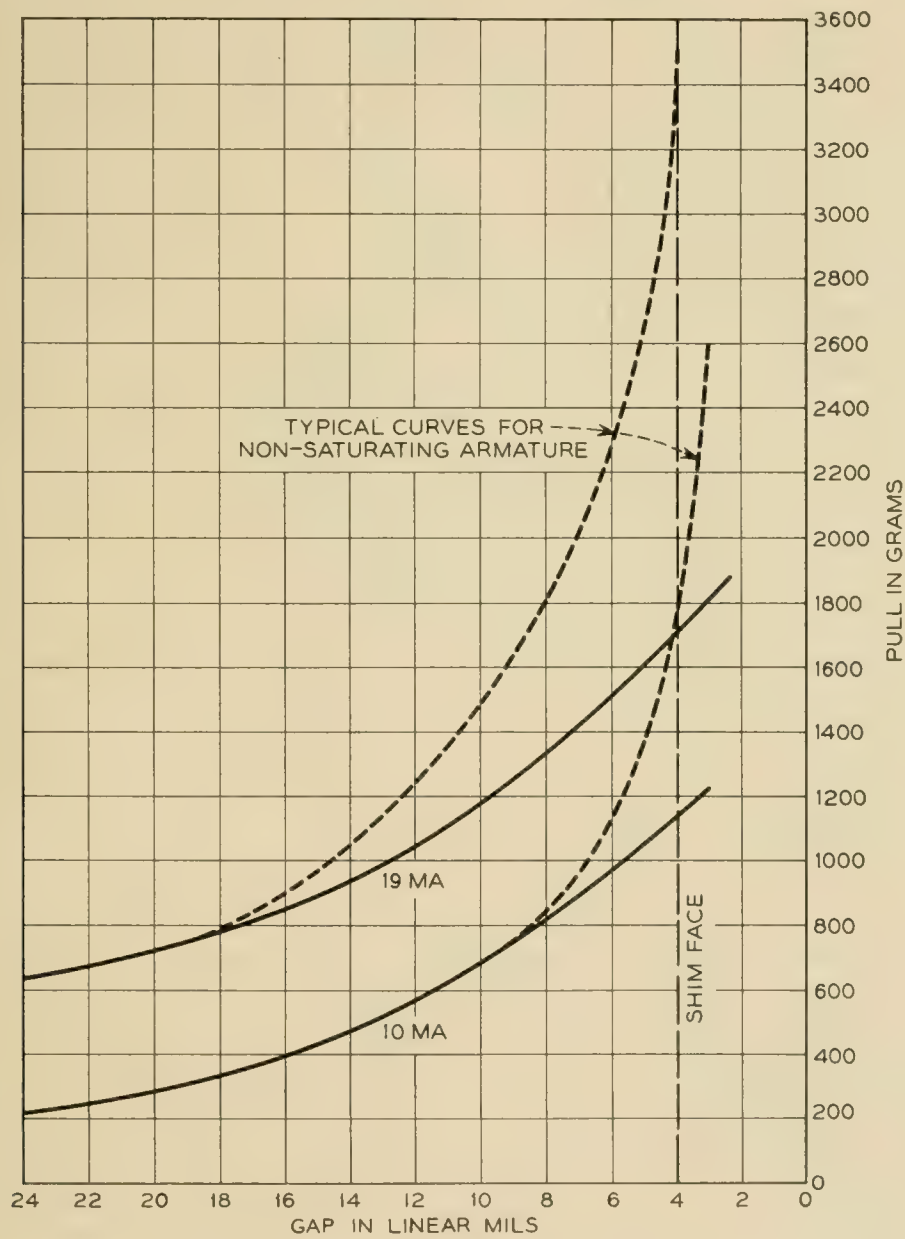


Fig. 9 — Magnetic pull curves for a coil of the J-7 valve.

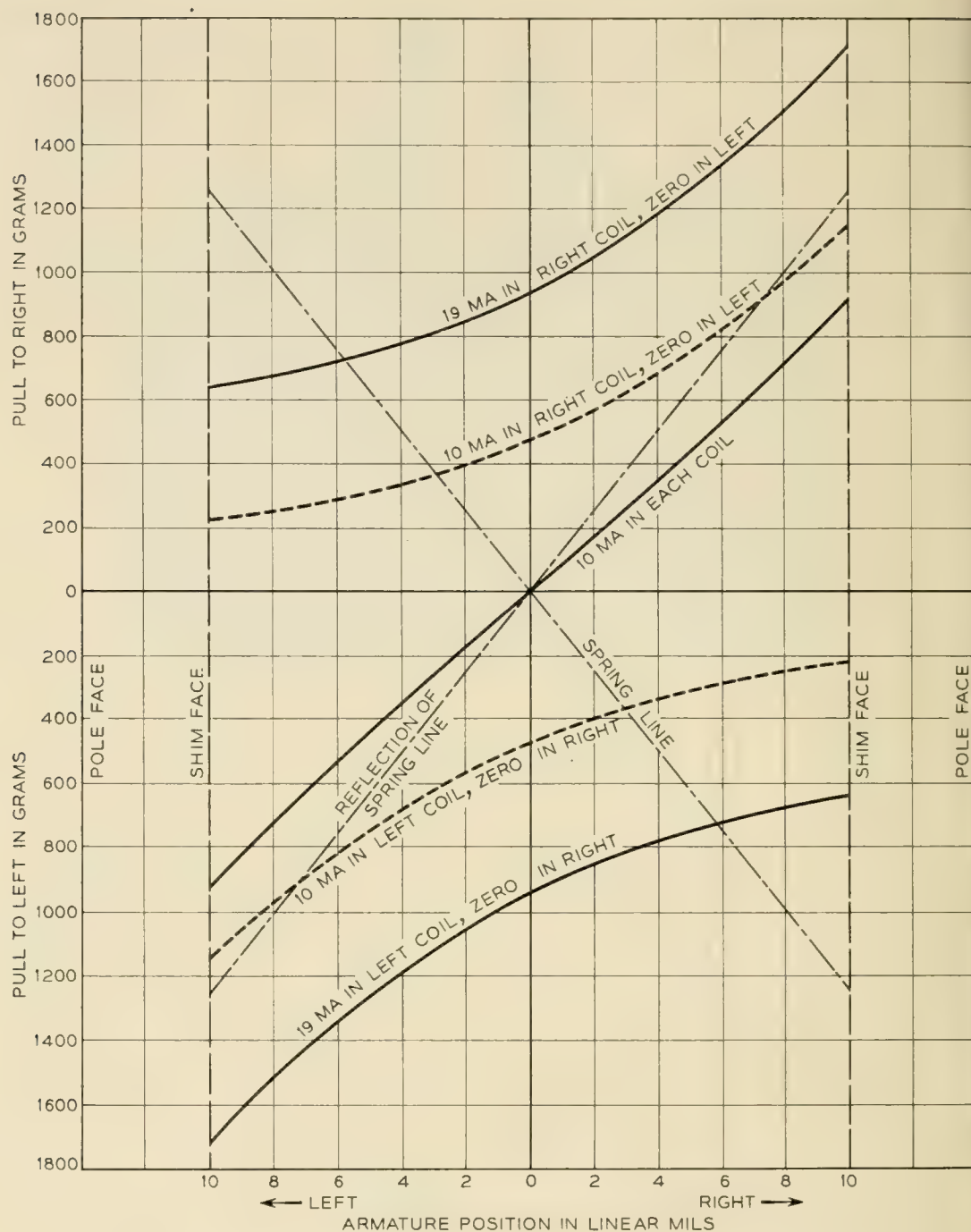


Fig. 10 — Spring and magnetic forces on armature of J-7 valve.

are shown to illustrate the need for these restrictive measures. The need is better appreciated when Fig. 10 is examined. This graph shows the differential pull of the two coils plotted against armature movement for extreme signals and for the balanced condition. There is also a line representing the spring force and one representing its reflection. The latter permits direct comparison of the magnitude of the opposing magnetic and spring forces.

It is important that the spring be able to center the valve when the coil currents are balanced. This means that the stiffness of the spring must be greater in magnitude than the negative stiffness created by the magnetic fields when 10 ma is flowing in each coil. On the other hand, the 19-ma current should be able to pull the armature against the pole piece and therefore must produce more force than the spring. If the shim and saturation limiting were not used it would be impossible to find a straight spring line that would fulfill both these requirements; i.e., its reflection would be between these curves without crossing either of them. The family of curves representing the net magnetic forces for the various intermediate values of current unbalance fall between the extreme cases shown. The intersection of one of these curves and the reflection of the spring line is the position which the armature will assume for that particular coil current. The reason for the relatively large margin of force shown for the 19-ma wide-open condition will be explained later.

Fig. 11 is plotted to show the net forces on the armature. The curves are the difference between the spring line and the two magnetic pull curves of Fig. 10. It can be seen that the forces are such as to cause the armature to move to the center in the balanced condition. In the case of maximum signal, it will move all the way to the shim stop in the direction of the coil which is carrying the current.

When there is no magnetic field present the armature resonance is about 320 cps. When measured statically, or at very low frequencies, with the coils energized, the negative stiffness of the field greatly reduces the effective stiffness of the spring, as seen in Fig. 11. However, when the valve is driven experimentally to find resonance, it occurs near 320 cps. This apparent increase in stiffness with frequency results from eddy currents that retard the change in flux to the extent that the negative stiffness virtually disappears. Eddy currents reduce the effective inductance of the coils from about 40 henries at very low frequency to less than 10 henries at 600 cps.

It is difficult to locate the resonance experimentally because of the large amount of damping provided by the extremely thin oil film between the plunger and the inserts. High resonance frequency of the valve is desired so that it is safely above any frequency encountered in the servo operation, thereby eliminating one consideration in the equalization. Also, a high resonance means that missile acceleration along the valve axis causes little displacement of the unbalanced mass of the armature.

A 250 cps differential dither voltage is superimposed on the push-pull dc signal to overcome the effects of static friction. The resulting 1 ma differential current produces a magnetic force about equal to that re-

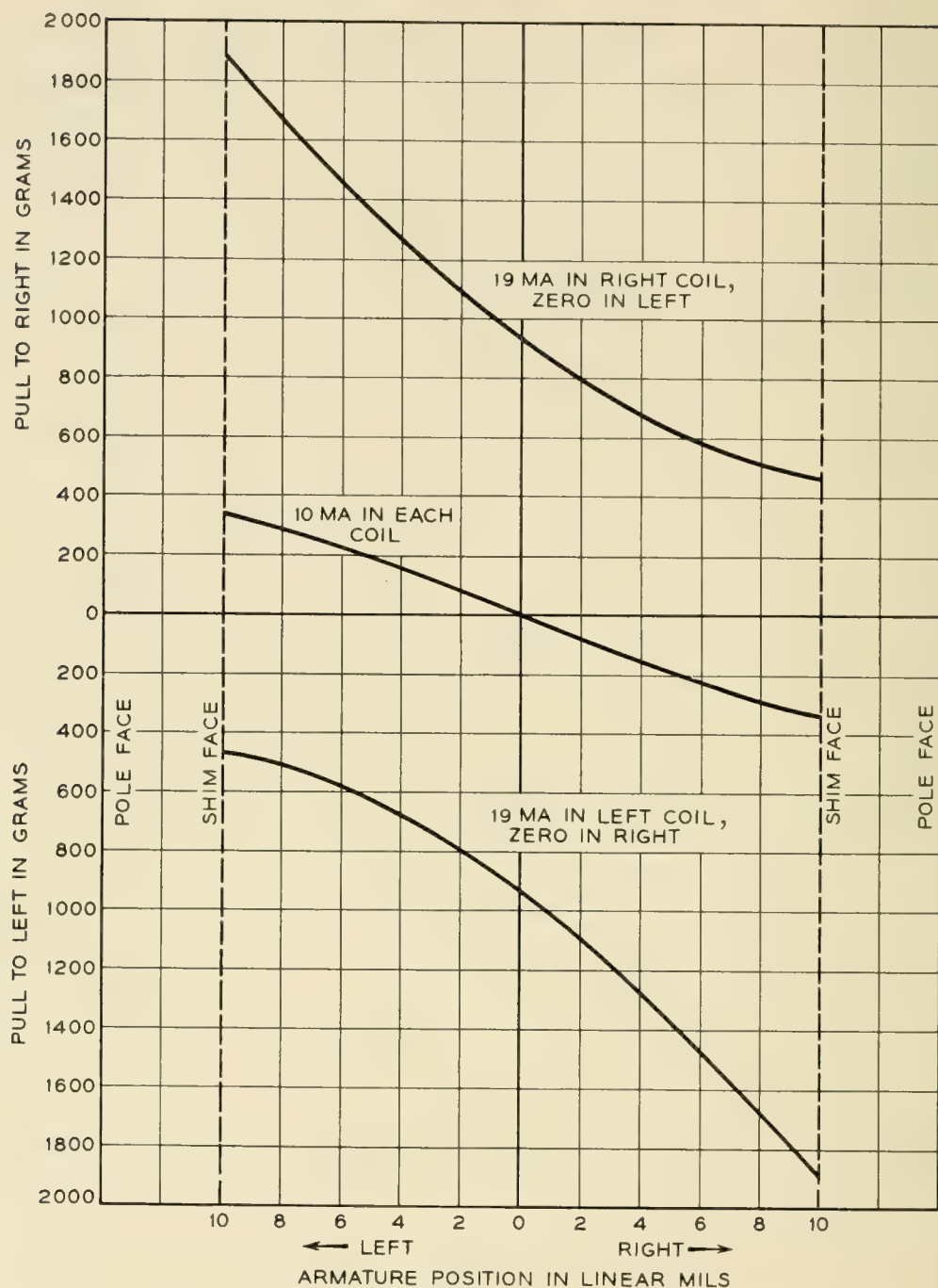


Fig. 11 — Resultant forces on armature of J-7 valve.

quired to move the armature in breaking the friction of the stationary armature assembly. Thus, the signal threshold is reduced by the amount of the dither current and the resulting increase in sensitivity to small signals greatly reduces the phase lags at low amplitude.

STEADY STATE HYDRAULIC CHARACTERISTICS

The J-7 valve is designed to operate from an oil supply having a pressure of 2,000 psi, which is somewhat higher than early valves of this

type. The increase in working pressure is a great advantage for a guided missile application because it results in lower weight, higher gain, and faster response. For example, doubling the pressure permits actuating cylinders of one half the size, an oil reservoir of one half the volume, and an increase in both response and gain by a factor of about three. Such features are sufficiently attractive to be worth a great deal of development effort. However, reliable and stable operation can be achieved under these high-gain conditions only if parasitic forces are kept extremely small.

It was found that the relation between pressure drop and flow is not so simple as one might expect from a sharp-edged, orifice-type control. For large openings of the control orifice, the pressure losses in the fixed orifices and passages of the valve body become an important factor. The following law is an adequate representation of pressure-flow characteristics:

$$p = 10q + \left(2 + \frac{455}{x^2}\right) q^2 \quad (1)$$

where

p = pressure drop, psi

q = rate of flow, cu in/sec

x = valve opening, linear mils

This equation was derived from test data from a model valve. These data confirmed a computational analysis of the hydraulic circuit. Fig. 12 graphically illustrates the equation. It is a plot of flow against valve position for various pressure drops. It will be noted that there is 0.001 inch difference between valve position and valve opening because of this amount of overlap at the ports.

Equation (1) provides the information necessary to compute the maximum output power of the valve. If all the pressure drop were across the control orifices, all the pressure would be utilized to accelerate the oil at this point and only the square term of (1) would exist. If this were the case, the maximum output power would occur when the pressure drop across the valve was one-third of the supply pressure. The other two-thirds of the pressure would be used to produce work in the cylinder. If laminar flow existed throughout the valve the square term would drop out, leaving a linear equation. If this were the case, maximum power would occur with the total pressure equally divided between valve and load. When both the linear the square terms are present, maximum power will occur when the pressure drop across the valve is somewhere

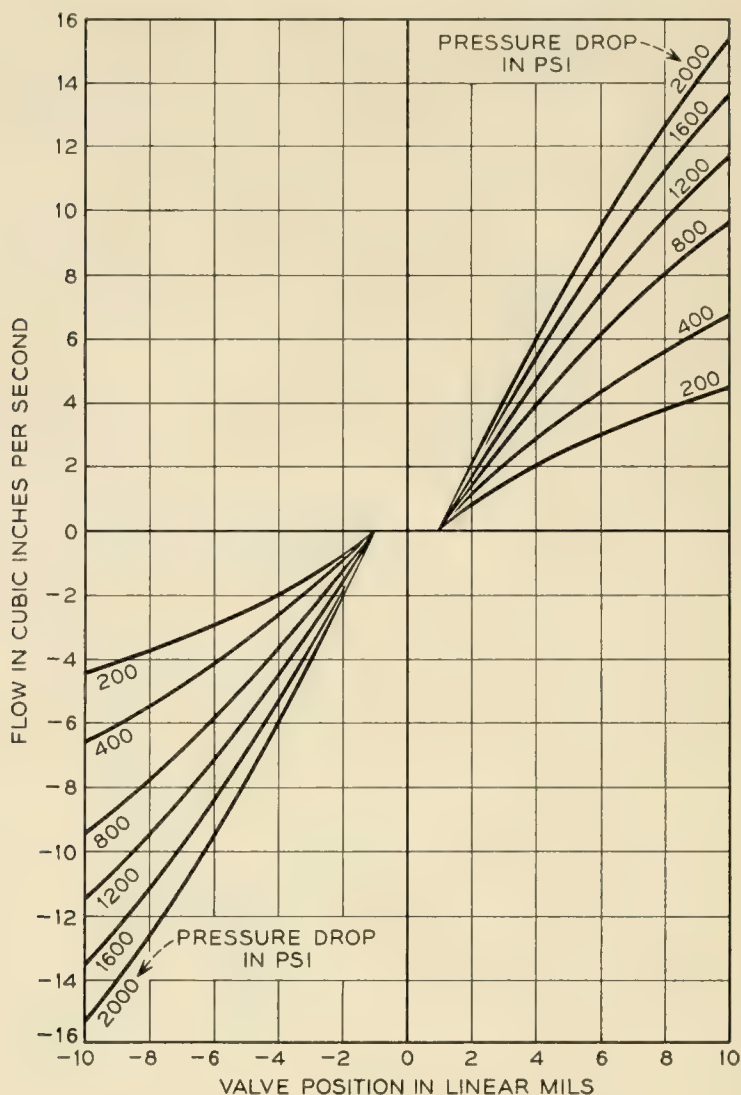


Fig. 12 — Flow characteristics of J-7 valve.

between one-third and one-half the supply pressure. The exact point is dependent on the magnitude of the supply pressure as well as the coefficients in the equation.

The valve will be wide open, an opening of 0.009 inch, when maximum power is produced. In this situation, (1) becomes:

$$p = 10q + 7.6q^2 \quad (2)$$

Useful output power at the load is the product of the flow rate and the pressure exerted on the piston.

$$\dot{W} = (P - p)q \quad (3)$$

where

$\dot{W}$ = Output power

P = Supply pressure, psi

Therefore

$$\dot{W} = Pq - 10q^2 - 7.6q^3$$

This expression can be differentiated and equated to zero to find the point of maximum power.

$$\frac{d\dot{W}}{dq} = P - 20q - 22.8q^2 = 0 \quad (4)$$

$$q = \sqrt{0.192 + 0.0439P} - 0.439$$

Substituting for q in equation (2) we find that for maximum power

$$p = 0.333 P + \sqrt{2.15 + 0.49P} - 1.46 \quad (5)$$

The normal supply pressure for the J-7 valve is 2,000 psi. Substituting this value for P in (4) and (5)

$$q = 8.95 \text{ cu in/sec}$$

and

$$p = 696 \text{ psi}$$

When these values are substituted in (3), we find

$$\begin{aligned} \dot{W} \text{ max} &= 11,700 \text{ in lb/sec} \\ &= 1.77 \text{ horsepower} \end{aligned}$$

Examination of the above equations will show that if the valve is used with a very low supply pressure, the linear term in (1) is dominant. In this case the maximum power output occurs when the pressure drop is nearly one-half the supply pressure. In the case of a very high supply pressure, the squared term is dominant and maximum power occurs when the pressure drop across the valve approaches one-third the supply pressure.

The ratio between the electrical quiescent input power to the coils and the maximum hydraulic out power is about 1,600, or a power gain of 32 db. Based on maximum signal the gain is 29 db. All forces must be precisely balanced and tolerances on parts be carefully controlled in order to realize this amount of gain in a single stage mechanical device.

DYNAMIC HYDRAULIC EFFECTS

Examination of the illustrations of the valve will show that it is statically balanced; i.e., pressure on any of the ports does not tend to translate the plunger. However, the flow of oil through a valve of this type produces a force on the plunger which tends to close the ports or

center the armature. This force creates a stiffness that adds to the effect of the mechanical centering spring. The magnitude of the force is quite nonlinear, it varies with pressure drop, and hence, with load as well as with armature displacement. Such variations tend to upset the stability of the servo loop in which the valve is used.

Fig. 13 is a simplified drawing of a valve which can be used to explain the dynamic effect. It will be noted that when the valve plunger is displaced in either direction, the fluid flow is metered by two control orifices in series. The oil flows from oil supply, through the valve body, into a groove in the plunger via the first of the two orifices. The second

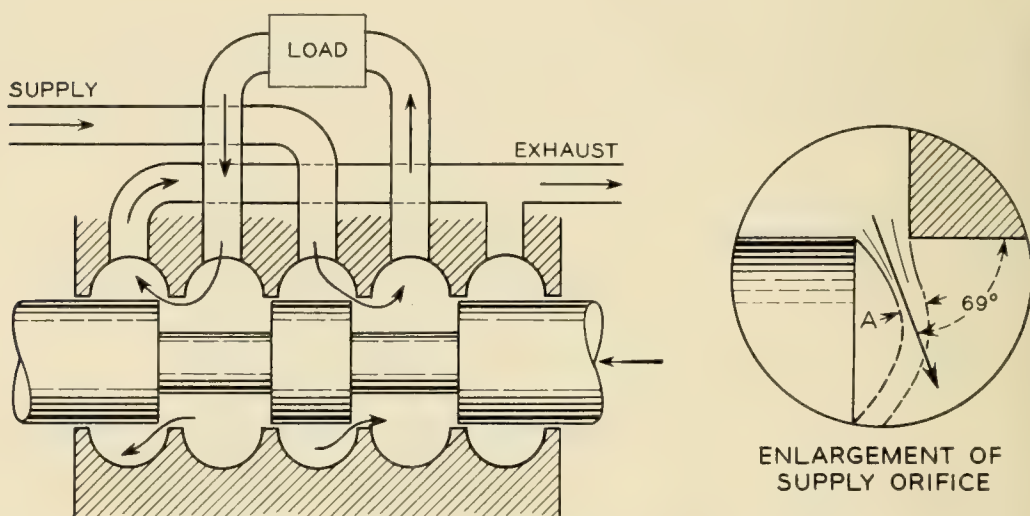


Fig. 13 — Simple valve with rectangular lands.

orifice meters the flow from the other groove in the plunger into the exhaust part of the valve body. Most of the pressure drop in the valve appears across the two orifices and is equally divided between them. The maximum fluid velocity occurs immediately downstream from the orifices at the vena contracta of the jet. This point is labeled "A" in the enlarged insert on Fig. 13. The velocity at this point can be computed by use of Bernoulli's Theorem. In the case of the J-7 valve, where the valve opening is small compared to other passage dimensions, many of the terms of the equation that formulates this theorem can be neglected. In this way the equation becomes

$$h = \frac{V^2}{2g} \quad (6)$$

$$V = \sqrt{2gh} \quad (7)$$

where

V = fluid velocity at the vena contracta, ft/sec

h = pressure drop across orifice, feet

g = acceleration of gravity, ft/sec²

Von Mises² has shown that the departure angle of the jet from a small orifice, such as in the configuration shown in Fig. 13, is 69° from the longitudinal axis. Tests made with an orifice, shaped like those in a simple valve, showed that the jet continues at this angle for a short distance only. Further downstream the jet turns to hug the radial surface on the plunger. This action is depicted by the dotted lines in the insert on Fig. 13. Bernoulli's equation explains that pressure is exchanged for velocity. The low pressure within the jet stream pulls it toward the nearest wall of the cavity. The flow of oil over this surface reduces the pressure on that wall and unbalances the distribution of forces on the surface of the annular grooves in the plunger. The area of reduced pressure causes a net longitudinal force in the direction to center the plunger or close the valve.

Examination of the situation around the exhaust orifice shows that a similar action will occur, but will not result in a comparable force on the plunger. The area of high velocity lies along a surface in the valve body rather than acting on the plunger. The velocity upstream from the orifice is not localized and hence produces forces that are small with respect to those downstream. The fact that the exhaust port forces on the plunger are small compared to those of the intake was confirmed by tests.*

There is a small time lag between the opening of the ports and the dynamic centering force which is proportional to the rate of change of oil flow. This lag results from the fact that finite time is required to change the velocity of the oil mass in the system. At high frequencies, the delay results in a considerable phase lag between the plunger position and the dynamic force. This delay means that fluid velocity, and hence the force, is higher during the quarter cycle in which the valve is closing than it is during the quarter cycle in which the valve is opening.

* Considerable work has been conducted on valve theory and design elsewhere since the J-7 valve was developed. Reference 4 is an excellent example of a thorough analysis of valve dynamics with an approach to the problem from a different viewpoint. This reference reports on tests and theories which show the secondary forces from the exhaust and intake orifices to be equal. This is in direct contradiction to the experience with the J-7 valve and remains as an unresolved problem in the mind of the writer.

Therefore, more energy is exerted on accelerating the mass of the plunger during the closing operation than is absorbed in slowing the mass during the opening phase. If the net gain in momentum is larger than can be absorbed by the viscous damping of the oil film in the lapped fit, oscillation⁵ will occur. In the case of the J-7 this oscillation tended to occur at slightly over 400 cps.

The Bernoulli force described above was recognized and measured early in the development of this series of valves, but was considered unimportant due to the high stiffness and large damping inherent in the design. The experimental models and the early production models showed no indication of oscillation. Later in the production program, the lapped clearances were increased and high ambient temperatures at the missile test locations were encountered. These two factors combined to reduce the damping, due to the working fluid, to the point where hydraulic oscillation occurred. An external damper was added to alleviate this problem. The damper consisted of an aluminum piston closely fitted to an aluminum cylinder. A viscous fluid (polyisobutylene) between these two parts provided sufficient damping to stabilize operation. This fluid also has the advantage of a relatively small decrease in viscosity with temperature. This type of damper is illustrated on the valve in Fig. 3. (Subsequent improvement in the internal design of the valve reduced the dynamic effect to the point where the need for the external damper has been eliminated.)

Fig. 14 shows a compensated intake orifice configuration corresponding to the insert picture on Fig. 13. It represents the first attempt to balance the dynamic or Bernoulli force. The depth of the annular groove

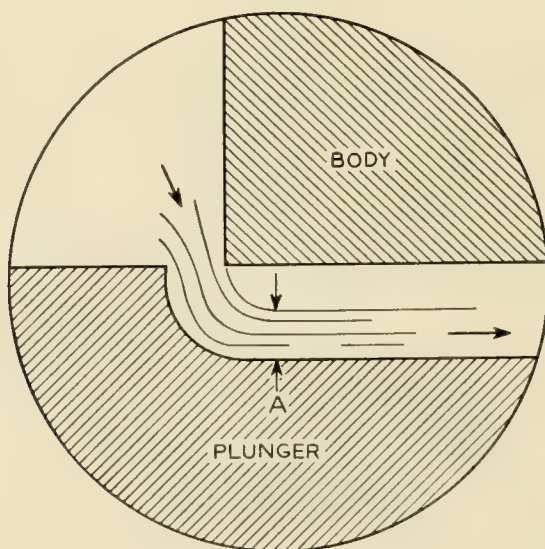


Fig. 14 — A compensated valve orifice.

in the plunger is reduced and a curved surface added to direct the flow parallel to the valve axis. This configuration reduces the dynamic force for two reasons. First, the amount of radial surface exposed to the low pressure is greatly reduced. Second, the curved surface acts much like a turbine blade in deflecting the oil stream and developing a reaction thrust that opposes the Bernoulli force. The reaction thrust increases as the longitudinal component of the high-velocity jet is increased. If the jet can be turned to become parallel with the longitudinal axis without appreciable loss in velocity, the maximum reaction thrust is obtained. In this case, the force is equal to the increase in the longitudinal component of momentum over the conditions of the free jet as shown in Fig. 13.

$$F = \frac{\rho q V}{g} (1 - \cos 69^\circ) \quad (8)$$

where

F = force, lb

ρ = fluid density, lb/cu in

q = flow rate, cu in/sec

V = fluid velocity, ft/sec

g = acceleration of gravity, ft/sec²

Calculations of the dynamic forces in accordance with the above reasoning yield only approximate results because the local velocities and their gradients are functions of passage shape as well as pressure drops. The contour of the plunger grooves were computed for use in the first experimental model, whose design was intended to alleviate this problem. Refinements to the initial model were made by cut-and-try methods. Since the forces involved are relatively small, and their magnitude changes so rapidly with plunger position, specialized measuring instruments had to be developed whose sensitivity was high and compliance very low.

A certain amount of contradiction was apparent in the force measurements made. To better understand the action of the oil within the valve, a transparent replica of the cross section of a valve port was used under a microscope. Figs. 15, 16, and 17 are illustrations of typical tests. Fig. 15 is two views of an early type valve with rectangular ports at different openings and pressure drops. The arrows indicate a portion of the cylindrical sliding surface separating the plunger and body. The lower left shadow is the sharp corner at the edge of the annulus in the plunger.

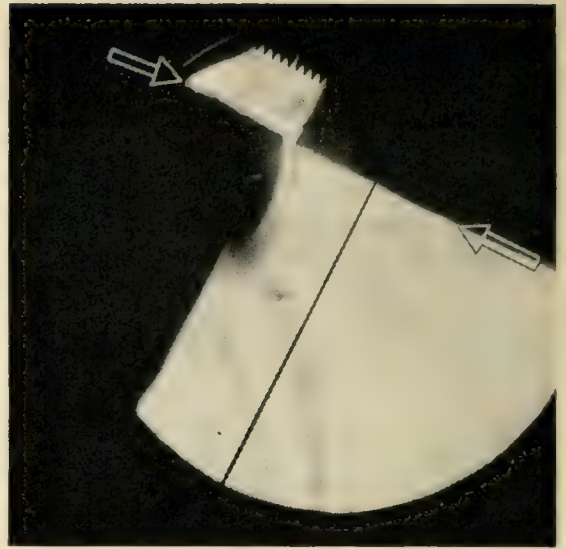
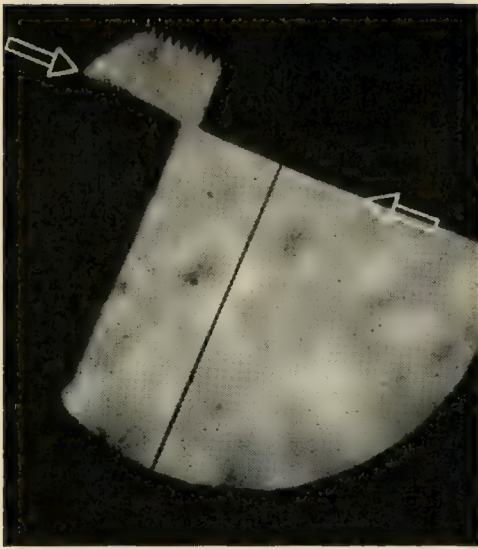


Fig. 15 — Oil flow through simple port.

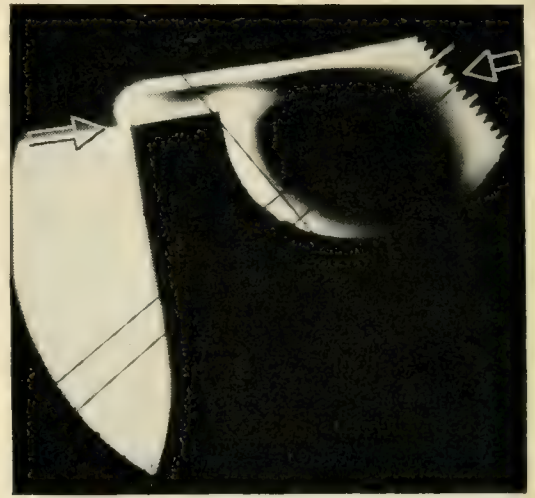
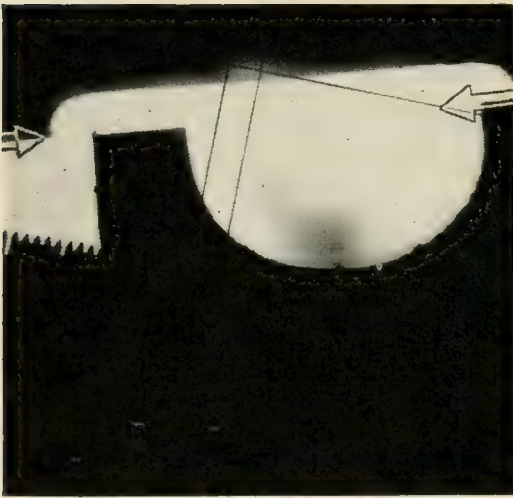


Fig. 16 — Flow through port with decreased Bernoulli effect.



Fig. 17 — Flow through port with decreased Bernoulli effect and increased flow.

The lower light area is oil in the annular groove of the plunger. The oil is flowing from top to bottom. It will be noted that the flow on the downstream side of the orifice hugs the radial surface of the plunger, as discussed above and depicted in Fig. 13. The crosshair and the comb-like scale are a part of the microscope.

Fig. 16 shows two views of a subsequent trial model quite similar to that shown in Fig. 14. Here, the bottom shadow is the insert and the top is the plunger. The flow is from bottom left to top right. This particular design reduced the Bernoulli force but resulted in a serious reduction in flow. This reduction was caused by the large cavitation bubble visible in the right view. This bubble was the result of rapid rotation of oil in the chamber. It prevented the orderly release of the oil through the internal passage to the actuating cylinder (not visible in pictures).

Fig. 17 shows a trial model similar to the J-7. The left illustration

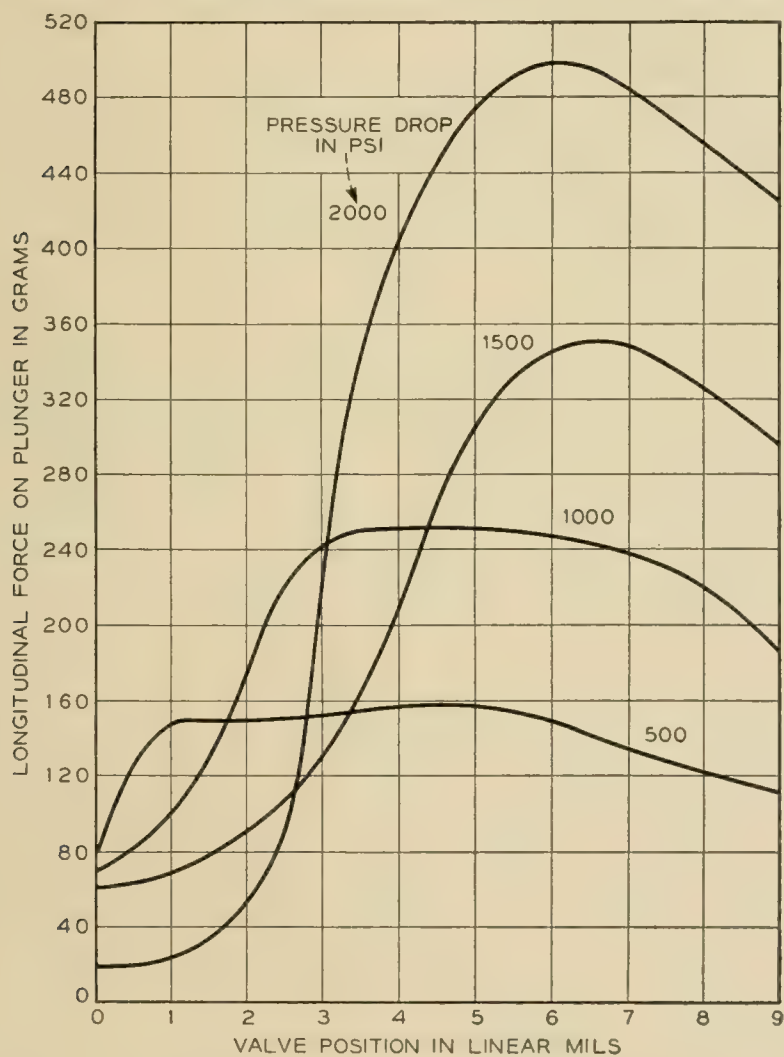


Fig. 18 — Force on plunger of J-7 valve due to Bernoulli effect.

shows the oil flowing from left to right and simulates the intake port of the valve. The extra land on the plunger directs the oil stream toward the escape passage to reduce turbulence and cavitation and increase oil flow. The right illustration illustrates the reverse flow of oil, right to left, and represents the exhaust port of the valve.

Fig. 18 is a plot of the measured Bernoulli force on the J-7 valve. It will be noted that there is little relation between the curves for various pressure drops. No simple equation has been formulated to account for the forces observed. Although Fig. 18 shows the Bernoulli force to be large, Fig. 11 shows that full signal produces enough net force on the armature to overcome this force at any valve position.

A large part of the development effort on the J-7 model was expended on the problem of reducing the forces caused by oil flow. For the same flow conditions, the J-7 has about one-fifth the Bernoulli force of earlier designs with simple rectangular grooves in the plunger.

Subsequent to the initial manufacture of the J-7 valve, the design of the annular grooves and body inserts has been improved continually. Consequently, the Bernoulli force has been further reduced to permit higher operating pressure, and hence more gain, without creating a hydraulic oscillation problem.

THE J-7 VALVE AS A SERVO ELEMENT

For any given set of operating conditions, the transfer function (expressed as cubic inches of oil flow per milliamper of control current unbalance per lb per square inch of pressure drop) can be extracted from the information presented above. However, the resulting family of curves for various load torques would be of little use to the servo designer. The pressure drop available for use by the valve is different for each curve, and they are all quite nonlinear, as is apparent from the data.

The overlap of the valve ports results in another type of nonlinearity that complicates the loop equalization problem. Examination of Fig. 12 will show that the effect of the overlap is a small dead area in the region of zero output of the valve. Small signal levels will cause the valve armature to operate in and around the vicinity of the dead zone, resulting in very little oil flow. Thus, the gain of the valve for very small signals is lower than for signals of greater magnitude.

As mentioned earlier, the effect of the nonlinearities of the valve is greatly reduced by use of the relatively fast-acting local loop which encompasses the valve, actuating cylinder, and amplifier. This inner, or secondary, loop contains sufficient gain to insure that, in spite of the

nonlinearities, the fin position is controlled in strict accordance with the summation of the input signals applied to the amplifier.

The first design of the servo circuits was made by using the slopes and magnitudes of typical and extreme points of operation as obtained from Figs. 11 and 12. The design of the servo equalization networks was obtained by successive refinements made during actual tests of the complete servo systems. These tests were performed with the aid of rather elaborate simulators that subjected the systems to the conditions of actual flight.

The characteristics of the valve and the amplifier which drives it, as applied to a servomechanism, can best be illustrated by plotting gain and phase shift versus frequency in an open loop. Fig. 19 is such a graph. This information was gathered by applying an input signal to the servo amplifier from an oscillator. The valve was driven by the amplifier in the usual manner. The valve controlled the flow of oil to a piston which operated a load that was equivalent to a typical aerodynamic load as seen by the control surface. The voltage from the fin-position potentiometer was compared to the amplifier input. Fig. 19 shows the phase and amplitude comparison of these two voltages. A small amount of feedback was used to prevent the piston from drifting to one end of the cylinder.

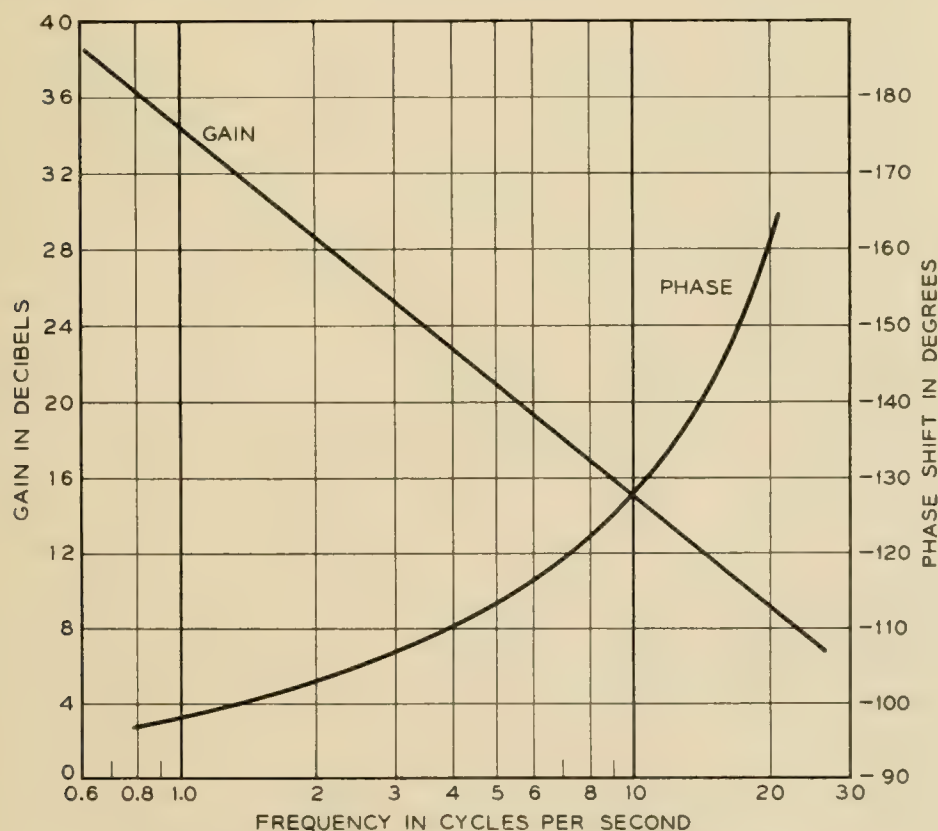


Fig. 19 — Frequency characteristics of J-7 valve.

The data were adjusted to correct for the error introduced in this manner, so that open loop conditions are represented. The equalization networks have compensating leading characteristics to prevent the oscillation suggested by the increasing phase shift at the higher frequencies.

If the servo acted in strict accordance with minimum-phase network theory, the slope of the gain curve would have started to increase at the high frequency end of Fig. 19. The effects of nonlinearities caused by such things as overlap and dither action cause this departure from classical theory. Actually, the slope does start toward 12 db per octave just above the frequency range shown.

Part of the phase shift shown on Fig. 19 is due to the inductance of the coils and the relatively low source impedance of the amplifier used. Later amplifier designs have much higher output impedance, which has greatly reduced this effect.

CONCLUSION

The series of hydraulic valves developed for use in the NIKE missile provide a light-weight, high-performance control element for positioning aerodynamic control surfaces. Although these electrohydraulic transducers provide a high power amplification, they are relatively simple single-stage devices. Their successful application in the NIKE missile has caused hydraulic servos to be considered for many other military control systems, some of which are under active development at this time. The hydraulic servo has many advantages to offer in the high power field that cannot be provided by other conventional types. It is expected that these advantages will foster a great increase in the use of hydraulic servos in high performance applications in the next few years.

ACKNOWLEDGMENT

The writer wishes to thank E. L. Norton who contributed greatly to all phases of the valve development and V. F. Simonick of the Douglas Aircraft Company who provided valuable consultation on the hydraulic problems involved.

REFERENCES

1. Peek, R. L. and Wagar, H. N., *Magnetic Design of Relays*, B.S.T.J., Jan., 1954.
2. Ekelöf, Stig, *Magnetic Circuit of Telephone Relays — A Study of Telephone Relays — I*, Ericsson Technics, **9**, No. 1, pp. 51-82, 1953.
3. Von Mises, R., *Berechnung von Ausfluss — und Ueberfallzahlen*, Zeitschrift des Vereines Deutscher Ingenieure, **61**, pp. 447-452, 469-474, 493-498, May-June, 1917.
4. Lee, Shih-Ying, and Blackburn, John F., *Axial Forces on Control-Valve Pistons*, Meteor Report No. 65, Mass. Inst. of Tech., June, 1950.
5. Holoubeck, F., *Free Oscillations of Valve-Controlled Hydraulic Servos*, Royal Aircraft Establishment, Technical Note Mechanical Engineering-100, Nov., 1951.

Strength Requirements for Round Conduit

By G. F. WEISSMANN

(Manuscript received October 16, 1956)

Underground conduits are subjected to external loads caused by the weight of the backfill material and by loads applied at the surface of the fill. These external loads will produce circumferential bending moments in the conduit wall. The magnitude and distribution of the bending moments have been determined by measurements of the circumferential fiber strains in thin-walled metal tubes subjected to the external loads. The effects of different backfill materials, different trench width, and trench depth have been investigated. Bending moments caused by static and dynamic loads have been compared. The bending moments are finally expressed in terms of the required crushing strength.

INTRODUCTION

For many years, vitrified clay has been the principal material for underground conduit used as cable duct by the Bell System. Vitrified clay conduit has, in general, given excellent service. It has more than adequate strength and durability for the wide variety of conditions under which it must be used and for the long service life expected of it. For this reason, relatively little attention has been given to the formulation of special strength requirements for this type of conduit during the period in which it has been standard for Bell System use.

For some time other types of conduit, mainly in the form of single duct, have appeared on the market. Many of their properties make them attractive enough to be considered for Bell System use. However, to prevent possible failure or excessive deformation of the conduit under field conditions each type of conduit should meet minimum strength requirements, in order to provide the same reliable service that clay conduit has given. The main purpose of this investigation was to determine the minimum strength requirement for round conduit under various field conditions.

An extensive investigation of the effects of external loads on closed conduits has been conducted at the Iowa Engineering Experiment Sta-

tion.^{1, 2} However, due to the large diameters of the conduits used and the particular test conditions employed, the test results obtained were not directly applicable for the determination of strength requirements for conduits for the Bell System.

Underground conduits under service conditions are subjected to external loads. These external loads are caused by:

- a. The weight of the backfill material.
- b. The loads applied at the surface of the fill.

The magnitude and distribution of the external loads around the conduit are affected by:

1. The properties of the backfill material.
2. The magnitude of the applied load at the surface of the fill.
3. The height of the backfill over the conduit.
4. The trench width.
5. The bedding condition.
6. The diameter of the conduit.
7. The flexibility of the conduit.
8. Impact.
9. Arrangements of conduits in the trench.
10. Consolidation and compaction of the backfill material.
11. Auxiliary protection of the conduit.

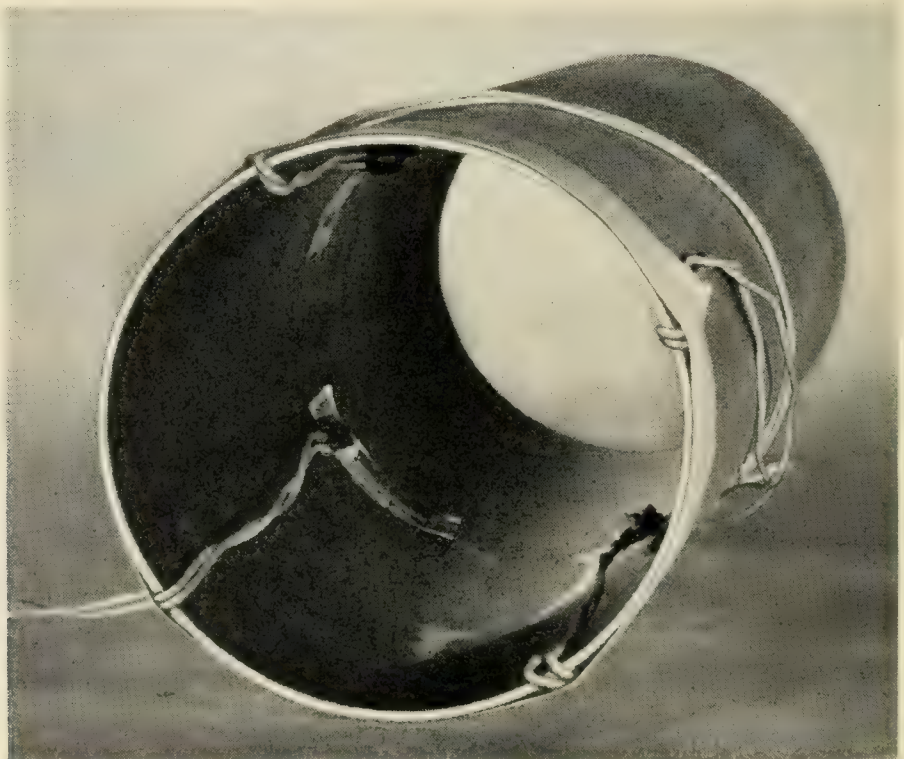


Fig. 1 — Thin-walled tube with SR-4 strain gages.



Fig. 2 — Test set.

External loads acting upon the conduit produce circumferential bending moments in the conduit wall. The magnitude and distribution of these bending moments have been determined in tests conducted recently at the Outside Plant Development Laboratory, Chester, New Jersey, and at Atlanta, Georgia. These tests were made with gravel, sand, and clay as backfill, in trenches of various width and depth and under conditions simulating, as nearly as possible, those encountered in the field.

TEST APPARATUS AND PROCEDURE

A test method was developed which permitted the determination of the circumferential bending moments in thin-walled conduits under field conditions.

The test device consisted of thin-walled aluminum or steel tubes of one foot length. The steel tubes used had an outside diameter of 4 inches and a wall thickness of 0.062 inches; the aluminum tubes had an outside diameter of 4.5 inches and a wall thickness of 0.065 inches. SR-4 strain gages were attached to the inside surface of the tube. Each tube was equipped with four equispaced SR-4 strain gages (type A-5), (Fig. 1). One aluminum tube contained, in addition, sixteen SR-4 strain gages (type A-8), which were equally distributed around the internal circumference of the tube. By means of an SR-4 strain indicator and an Edin brush recorder it is possible to measure the strains caused by static, as well as by dynamic loads. The strains could be measured with an accuracy of $\pm 10 \times 10^{-6}$ inch per inch.

TABLE I — LIST OF TESTS AND TEST CONDITIONS

Undisturbed Soil	Backfill Material	Height of Cover	Trench Width	Applied Load	Test Tube
A — Field Tests					
Georgia Clay	Clay	(inches) 24, 30, 36, 42, 48	(inches) 18, 22, 30	(lbs) 2250-8500	Aluminum Tube
Georgia Clay	Fine Sand	24, 30, 36, 42, 48	18, 22, 30	2250-8500	Aluminum Tube
Georgia Clay	Clay	24, 30, 36	22	500 lb (dropped 5 and 3 feet)	Aluminum Tube
Georgia Clay	Fine Sand	24, 30, 36	22	500 lb (dropped 5 and 3 feet)	Aluminum Tube
Chester Soil	Clay	18, 24, 30, 36	5	1275-7775	Aluminum Tube
Chester Soil	Fine Sand	18, 24, 30, 36	5	1275-7775	Aluminum Tube
Chester Soil	Fine Sand	24, 30, 36, 40, 48	24	1000-10175	Steel Tube
Chester Soil	$\frac{3}{4}$ in. gravel	24, 30, 36, 40, 48	24	1000-10175	Steel Tube
B — Laboratory Tests					
1. Two-point load					
2. Investigation of bedding					

The test equipment used is shown in Fig. 2. Each tube was laid length-wise on the trench bottom between two pieces of plastic conduit of the same outside diameter. The tubes were oriented so that one of the strain gages was at the top of each tube. The trench was then filled with the backfill material. The loads at the surface of the fill were applied by using trucks with various measured wheel loads. The trucks were either moved slowly to a stop in position over each tube, or driven at moderate speed across the trench. Additional impact tests were made with a Hydrahammer, which consists of a 500-pound-weight dropped from different heights onto the surface of the fill. Each tube was subjected,

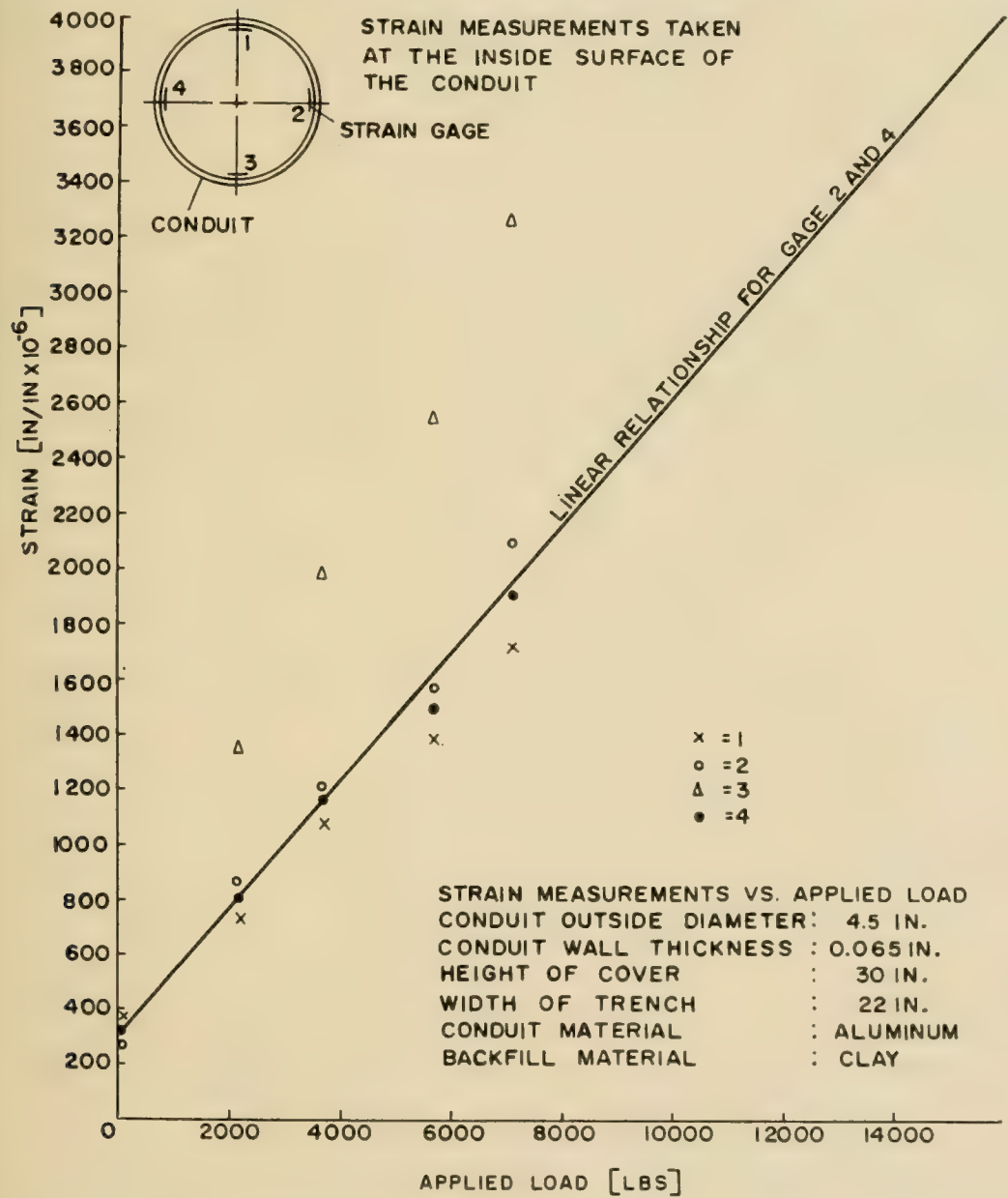


Fig. 3 — Strain measurements versus applied load.

in the laboratory, to various two-point loads (two edge bearing loads). Strain readings were taken during each test.

Table I lists the tests conducted and the test conditions investigated.

TEST RESULTS

Field and laboratory measurements obtained from each strain gage were plotted as a function of the applied load. A typical example for a field measurement is shown in Fig. 3. For this example, as well as for several hundreds of similar measurements, a linear relationship between the measured strains and the applied loads could be observed. For each case this linear relationship was derived from the data using the method of least squares. The straight line in Fig. 3 was plotted by this method and is shown with the data obtained in the test.

MOMENT DISTRIBUTION

Soil pressure acting upon a thin-walled tube will cause circumferential forces and bending moments in the wall of the tube. Furthermore, it is assumed that the strains caused by the compressive forces are small compared with those caused by the circumferential bending moments and are therefore neglected. For pure bending, the following relationship for circumferential bending moments and fibre strains is established.

$$M = \frac{1}{6} h^2 E \epsilon \quad (1)$$

where

- M = circumferential bending moment per unit length (in lb/in)
- h = wall thickness of tube (in)
- E = modulus of elasticity (psi)
- ϵ = circumferential fibre strains (in/in)

The circumferential bending moment of a thin-walled tube of unit length subjected to a two-point load is determined analytically:

$$M = \frac{Fr}{2} \left(\frac{2}{\pi} - \sin \theta \right) \quad (2)$$

where

- M = circumferential bending moment per unit length (in lb/in)
- F = applied two-point load (lb/in)
- r = radius of the tube (in)
- θ = angle with vertical axis of tube

The calculated circumferential bending moments (2) and those deter-

mined by strain measurements and Equation (1) are compared in Fig. 4. The agreement is very close.

Fig. 5 shows the theoretical moment distribution in a thin-walled tube subjected to a uniformly distributed vertical pressure; Fig. 6 is the theoretical moment distribution in the same tube subjected to a uniformly distributed vertical pressure at the top of the tube and a point load at the bottom of the tube.

Fig. 7 shows the typical experimental moment distribution in a thin-walled tube in an 18-inch wide and 40-inch deep trench with a backfill (36-inch cover) of moist Georgia clay subjected to an applied surface load of 10,000 lb. Fig. 8 shows the moment distribution under the same conditions but with a backfill of moist fine sand.

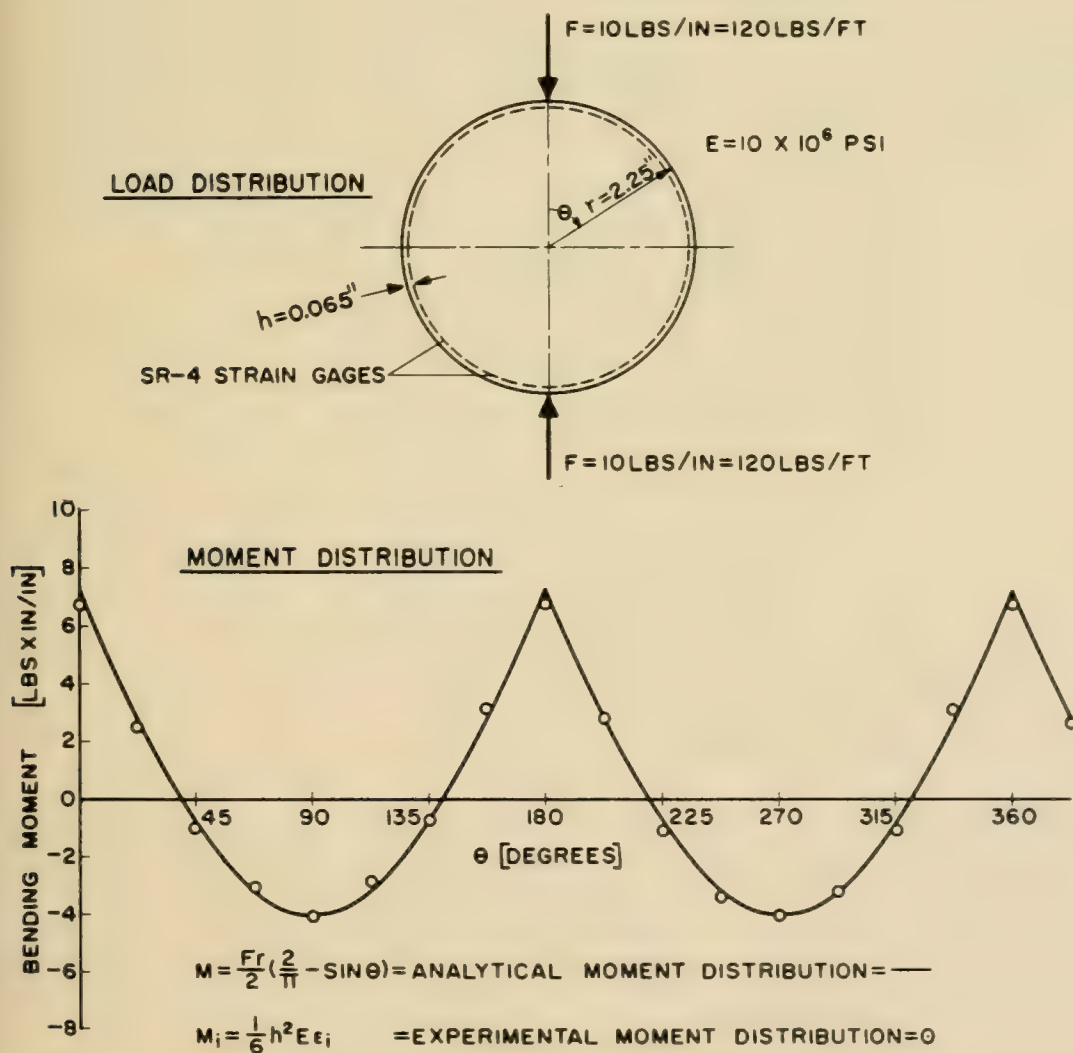


Fig. 4 — Analytical and experimental moment distribution in a thin-walled tube under two-point load.

BEDDING EFFECT

Based on the theoretical considerations illustrated in Figs. 5 and 6, comparison of Figs. 7 and 8 indicates a high load concentration at the bottom of the tube buried in a trench with clay backfill. For a verification of this assumption, a series of additional tests were conducted. The test tube was placed upon (a) a flat steel plate and then covered with moist clay, (b) a flat steel plate and then covered with dry sand, and (c) carefully distributed moist clay and then covered with clay. The external load was then applied. The moment distributions obtained for cases (a), (b), and (c) are shown in Figs. 9, 10 and 11, respectively.

A comparison of the data presented in Figs. 9 and 11 shows the effect of the bedding condition on the moment distribution (steel plate versus clay bedding). These data, as well as theoretical considerations (Figs. 5 and 6), indicate that the bedding condition affects mainly the bending moment at the bottom of the tube and only to a lesser degree the bending moment at right angles to the vertical axis. Due to a change in

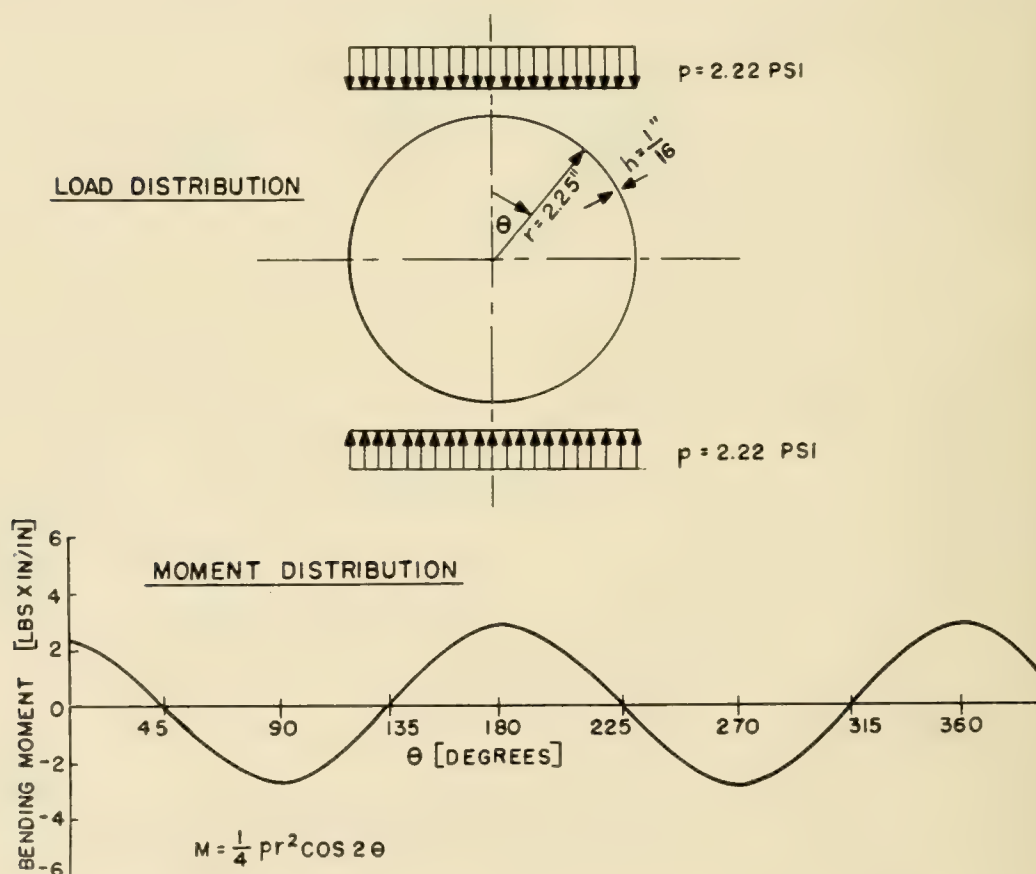


Fig. 5 — Theoretical moment distribution in a thin-walled tube subjected to uniformly distributed vertical pressure.

bedding, the moment at the bottom may vary up to 235 per cent, and the moments at the side points may change a maximum of 12 per cent.³ The field data indicate that bedding is of particular importance with moist clay backfill, and to a lesser degree for moist fine sand. However, bedding did not appear to affect the moment distribution of the tube for a dry sand or gravel backfill.

TRENCH WIDTH

The effect of the trench width on the magnitude of the bending moments in a conduit has been investigated. The presently available test results indicate the following:

a. No significant difference in the magnitude of the bending moments of the tube (4.5-inch diameter) could be observed for a trench width

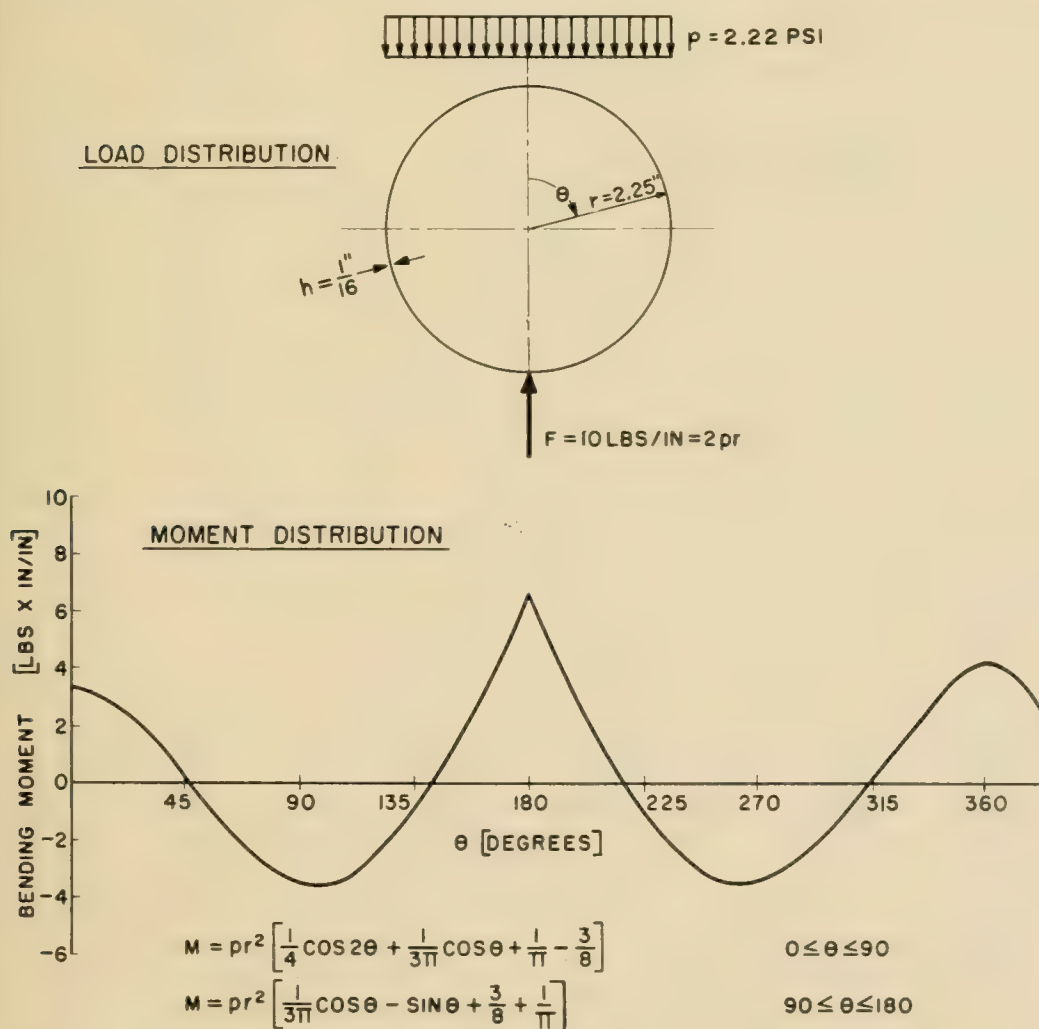


Fig. 6 — Theoretical moment distribution in a thin-walled tube subjected to uniformly distributed vertical pressure at the top and point load at the bottom.

between 18 and 30 inches. The magnitude of the bending moments appeared to be only a function of the backfill material, the applied load, and the height of backfill over the conduit.

b. For a trench width of 5 inches, the bending moments seemed to be independent of the type of backfill material. The same results have been obtained with a backfill of clay, sand, or gravel. This phenomenon indicates that the applied load was carried mainly by the undisturbed soil but deformation of the trench walls together with friction between the backfill material and the trench walls transmitted the load to the conduit. The magnitude of the bending moments in 5-inch wide trenches

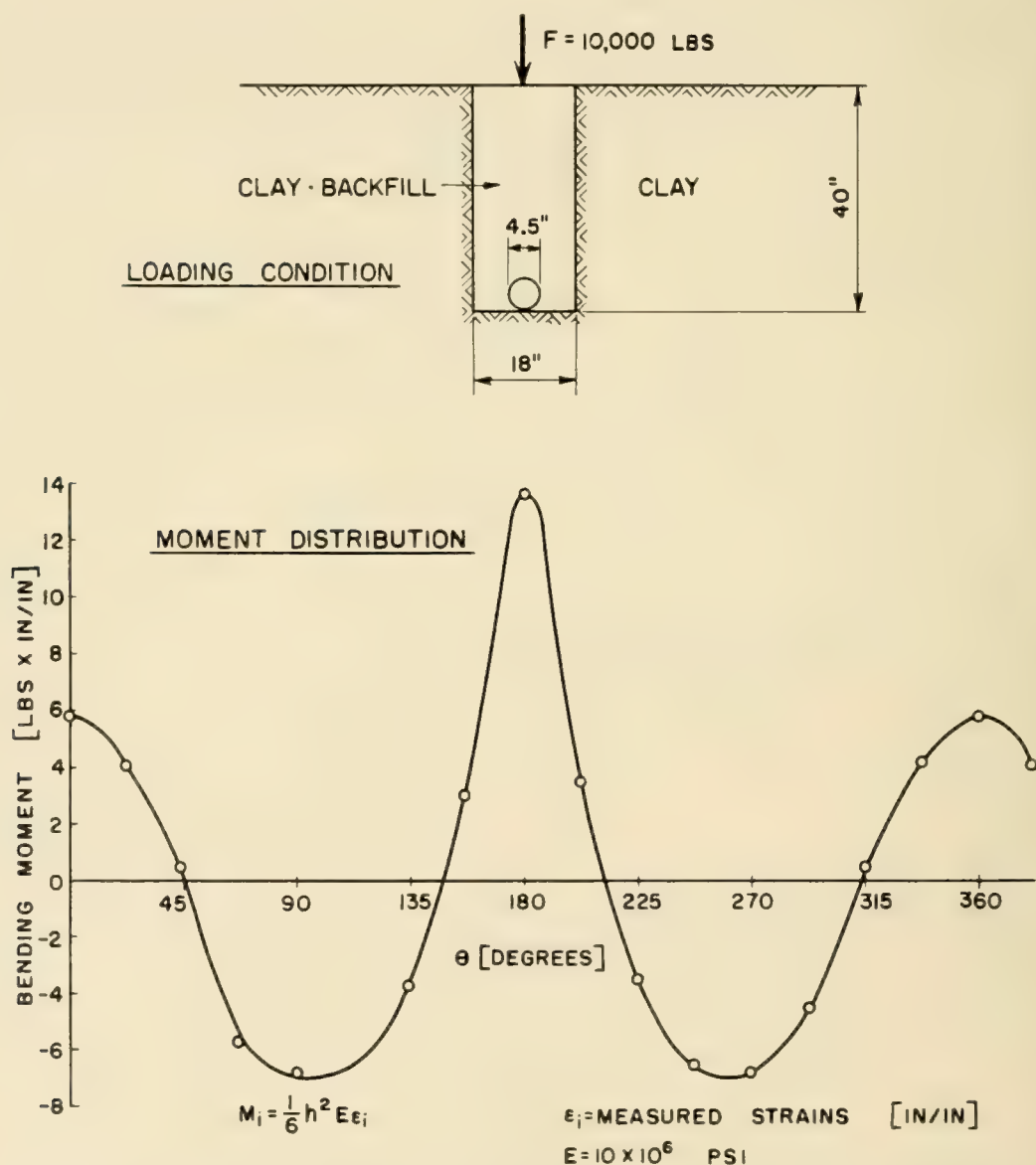


Fig. 7 — Moment distribution in a thin-walled tube, using clay backfill.

depends on the characteristics of the undisturbed soil. A comparison of the values obtained from a 5-inch wide trench dug in sandy clay (typical Chester, N. J. soil) and trenches between 18 and 30 inches wide shows that the values for 5-inch width are greater than for an 18- to 30-inch width when using sand backfill, and less when using clay backfill.

(Text continued on page 750)

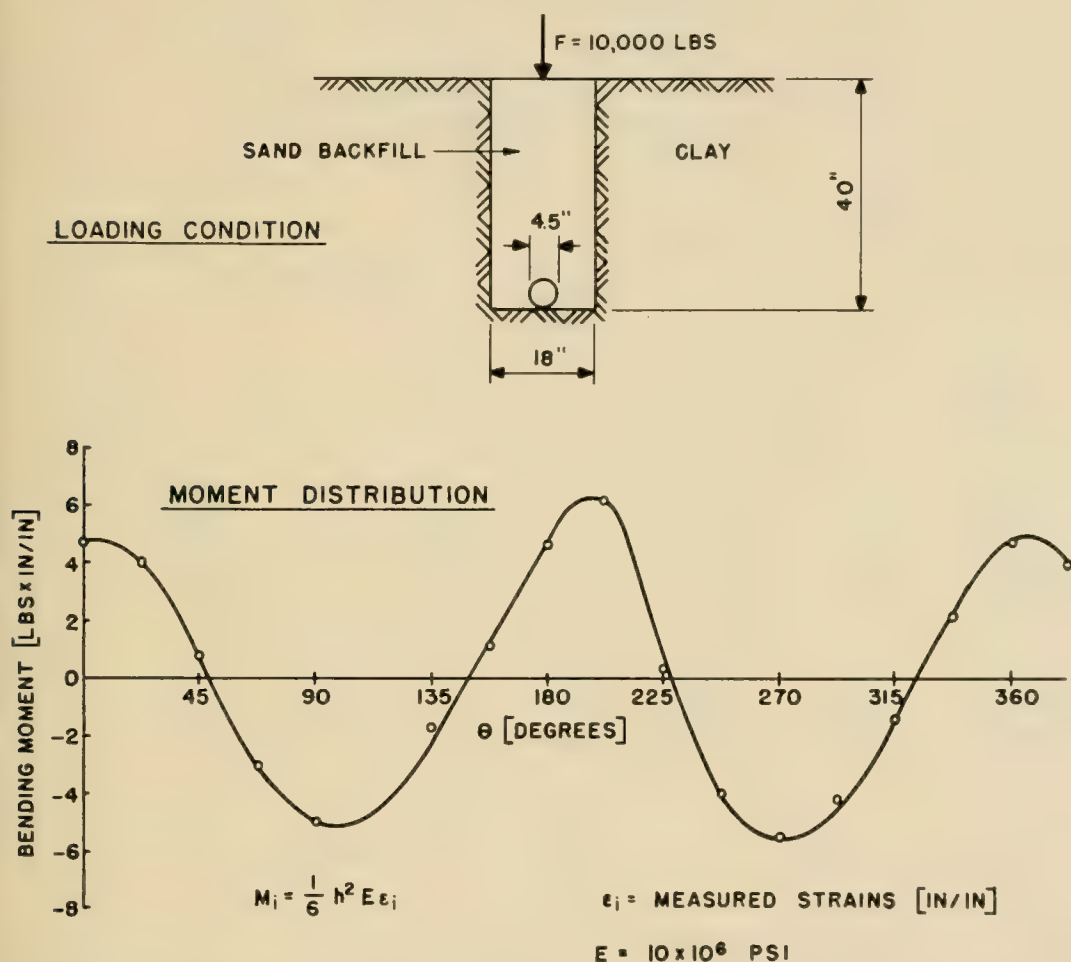


Fig. 8 — Moment distribution in a thin-walled tube, using sand backfill.

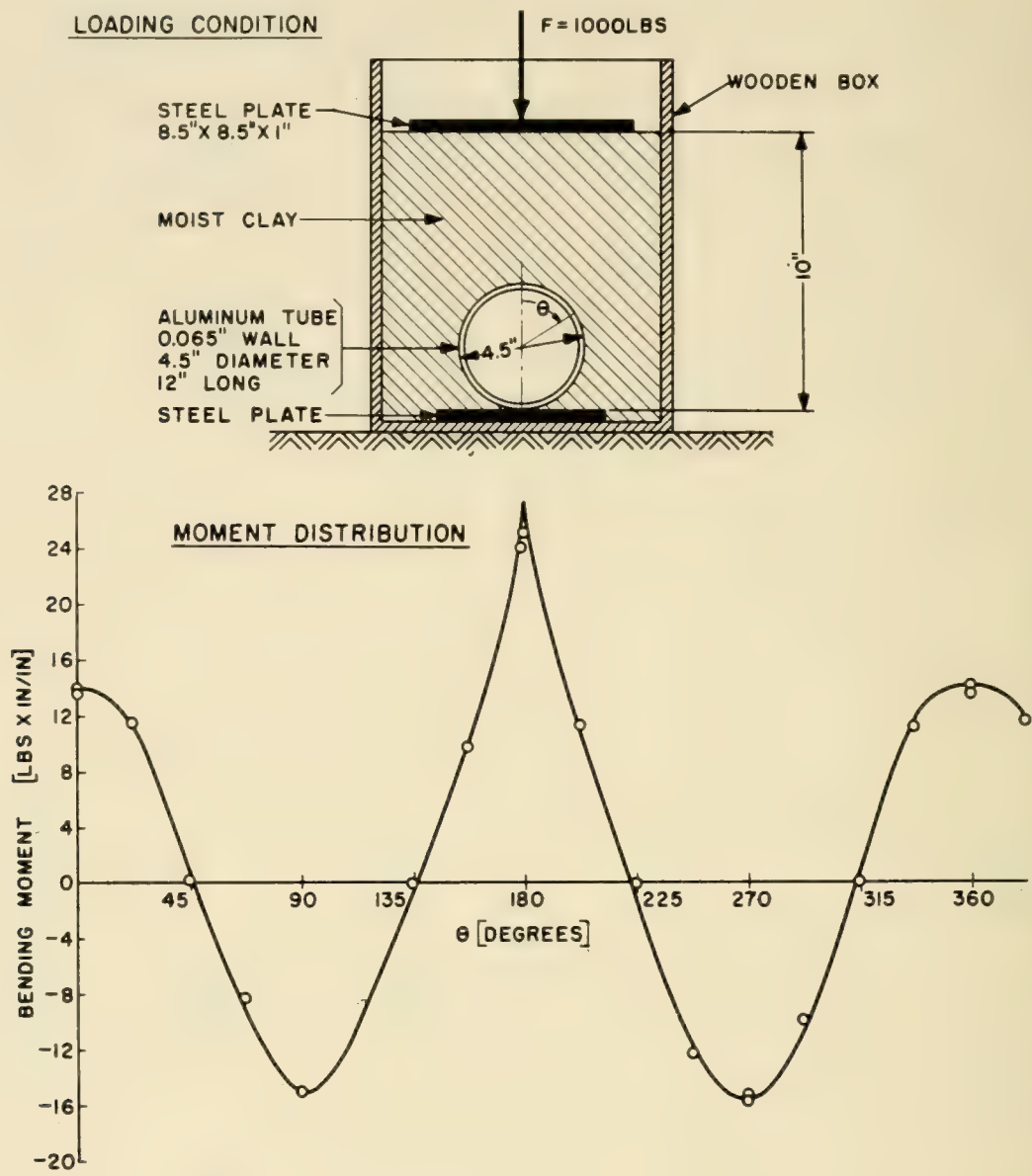


Fig. 9 — Moment distribution in thin- walled tube resting on steel plate and covered with moist clay.

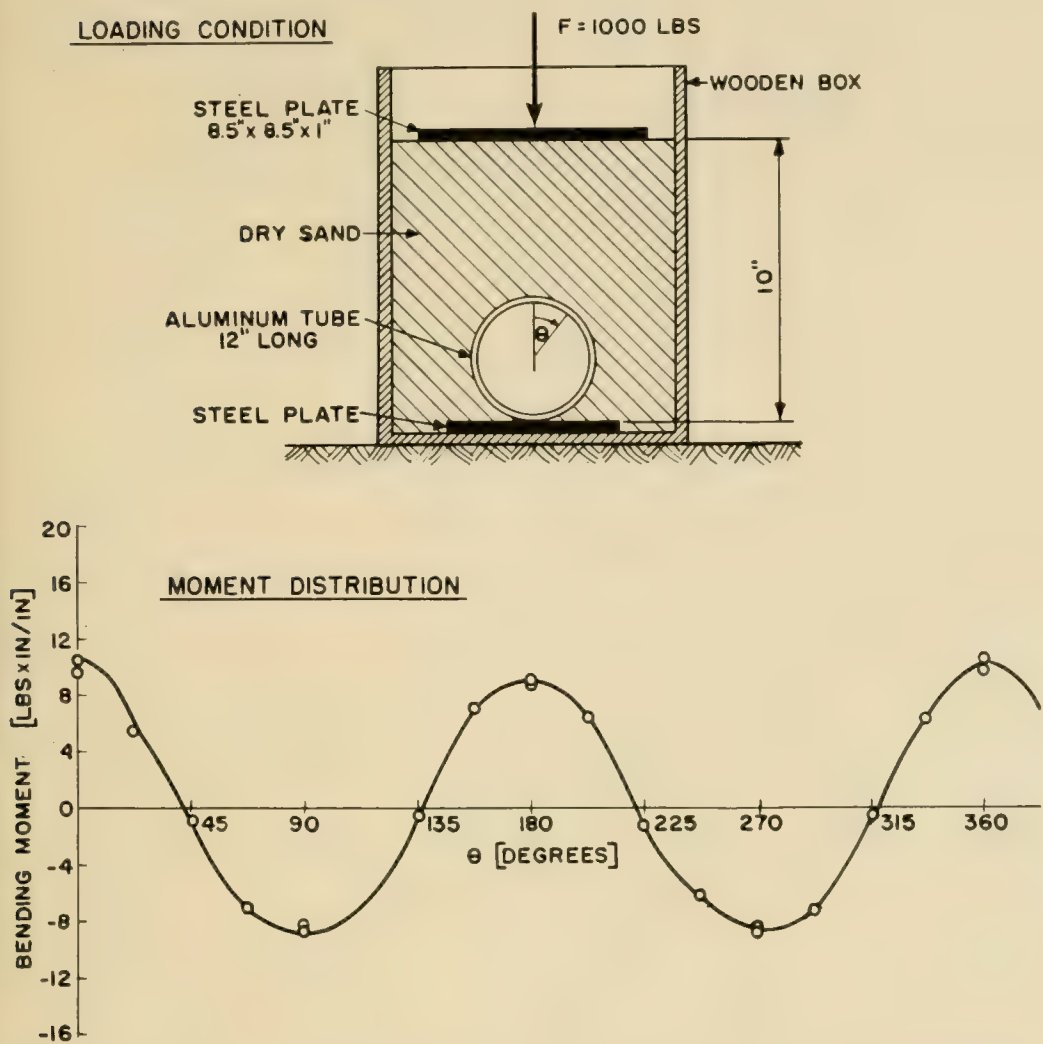


Fig. 10 — Moment distribution in thin-walled tube resting on steel plate and covered with dry sand.

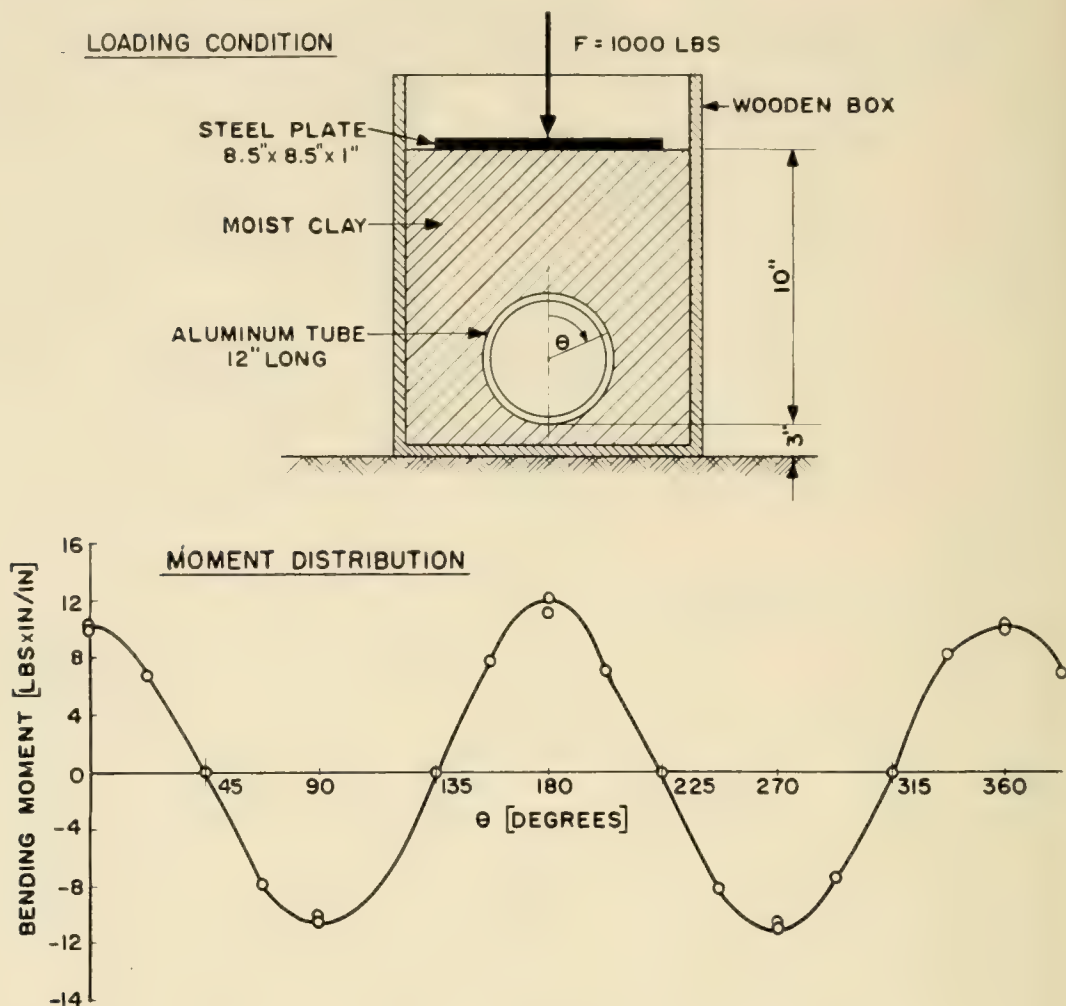


Fig. 11 — Moment distribution in a thin-walled tube in moist clay.

BACKFILL MATERIAL AND HEIGHT OF BACKFILL

Backfill provides the main protection for conduits against loads applied at the surface of the fill. The type of backfill used and the height of the backfill over the conduit are therefore extremely important. Figs. 12, 13, and 14 show the relationship between the bending moments in the conduit and the height of backfill for different applied loads. Fig. 12 shows this relationship for clay as backfill, Fig. 13 for fine sand, and Fig. 14 for gravel. The maximum bending moment is here expressed in terms of the "equivalent two-point load." The equivalent two-point load is the two-point load that will cause the same maximum bending moment in the conduit wall as the maximum bending moments measured in the field. The relationship between the two-point load and the bending moments in the walls of the conduit was given by (2). The equivalent two point load is obtained for $\theta = 0$ and becomes

$$F_{TP} = \frac{12\pi M_{\max}}{r} \tag{3}$$

where

- F_{TP} = equivalent two-point load (lb/ft)
- $M_{\max}$ = maximum bending moment in conduit (in $\times$ lb)
- r = radius of conduit (in)

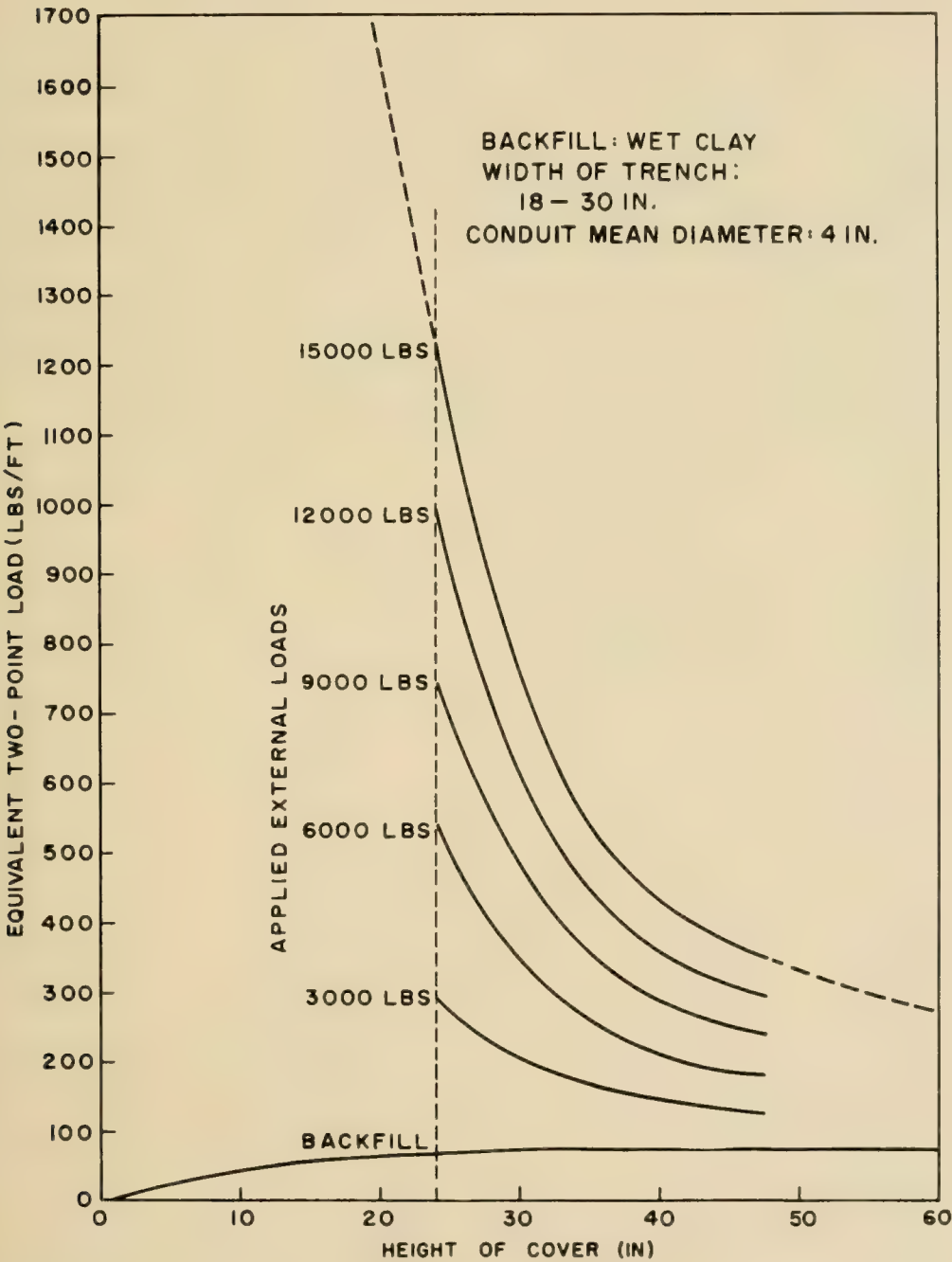


Fig. 12 — Equivalent two-point load versus height of cover for clay backfill.

The two-point load was chosen as a measure of the bending moment for convenience in laboratory testing of new conduits. The two-point load system is undoubtedly the simplest method for testing cylindrical or near-cylindrical shapes in conventional compression testing machines (crushing test). Since most of the supplementary conduit materials are cylindrical in shape, the most obvious requirement would be in terms of minimum two-point loading.

The data in Figs. 12, 13, and 14 were obtained as follows:

The bending moments at points perpendicular to the vertical axis were represented, as a function of the applied load, by a straight line, determined by the method of least squares. This was done for all the investigated depths and backfill materials. The side points were used as a basis of comparison because they were less affected by a change in the bedding conditions. With reference to the considerations of the

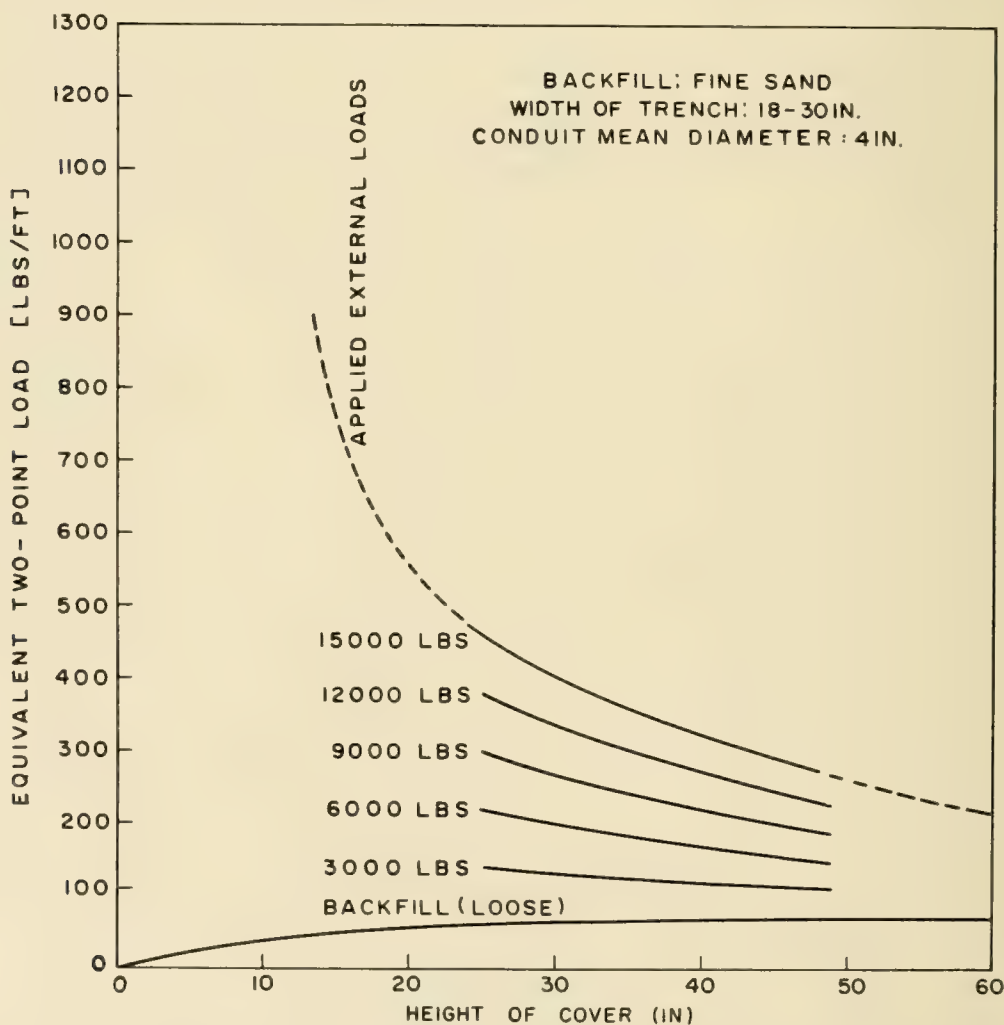


Fig. 13 — Equivalent two-point load versus height of cover for sand backfill.

bedding effect, the possible maximum bending moment at the bottom of the tube was obtained by doubling the values of the side points. Figs. 12 and 14 show clearly that the maximum bending moments occur in wet clay, and also that improvement is obtained by an increase in the height of backfill or a decrease of the applied load. These figures apply to 4-inch diameter tubes. In conversion of results obtained using 4.5-inch diameter tubes, it was assumed that the equivalent two-point load is directly proportional to the tube diameter.

DYNAMIC LOAD

A study was made to determine the effect of moving loads on the maximum bending moment of the conduit. For this purpose trucks were driven over the backfill at a speed of approximately 20 miles per hour, the strains measured and then compared with those obtained by static loads. The results show that for a clay backfill, the bending moments due to dynamic loads were equal to or even smaller than those obtained by static loads. For sand backfill, however, the dynamic loads caused an increase of the maximum bending moments of approximately 10 per cent. These results are in close agreement with dynamic load tests con-

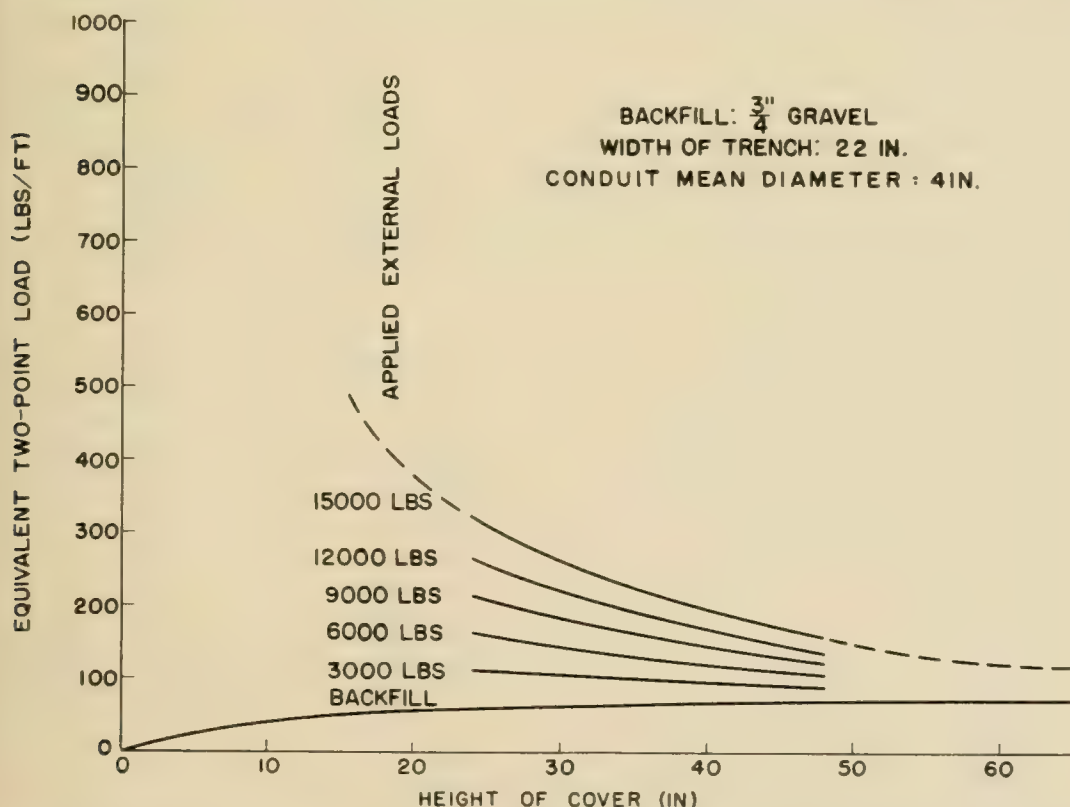


Fig. 14 — Equivalent two-point load versus height of cover for gravel backfill.

ducted by H. Lorenz.⁴ Further tests conducted with 500-lb weights dropped from different heights demonstrated that the effect of dynamic loads for clay backfill is of minor importance.

FURTHER INVESTIGATION

The tests previously conducted did not include investigations of the effects of conduit flexibility, various conduit diameters, multiple arrangements of the conduits in the trench, consolidation and compaction of the backfill, and auxiliary protection of the conduit. Studies of these factors are now in progress.

SUMMARY

Strength requirements of underground conduits depend on various factors. The most severe conditions were obtained with wet Georgia clay as backfill for a trench width of 18 to 30 inches. The relationship between the height of backfill, the load applied at the surface of the backfill, and the equivalent two-point load for these conditions is shown in Fig. 12. For example, a round conduit of 4-inch mean diameter under 24 inches of clay cover subjected to an applied load of 15,000 lb should be able to carry a two-point load of 1,250 lb without fracture or excessive deformation. Figs. 12 to 14 can be used to determine the required crushing strength of 4-inch round conduits subjected to different applied loads, buried at different depths with clay, sand, or gravel cover, but only for trench widths of 18 inches or more. Although this investigation is not completed, the results obtained to date may be used, within the indicated limits of application, for the selection of acceptable conduit.

ACKNOWLEDGEMENT

The author wishes to acknowledge with thanks the invaluable assistance given by R. G. Watling, W. T. Jervy, J. J. Pauer and Duncan M. Mitchel in planning and conducting these tests; also to express his gratitude for advice given by I. V. Williams during the progress of the investigation and preparation of this paper.

REFERENCES

1. Marston, A., The Theory of External Loads on Closed Conduits in the Light of the Latest Experiments, Bulletin 96, Iowa Engineering Experiment Station.
2. Spangler, M. G., Soil Engineering, International Textbook Company, Scranton, 1951.
3. Marquard, Rohrleitungen und geschlossene Kanäle, Handbuch für Eisenbetonbau, 4, Aufl., Bd. 9, Berlin 1939.
4. Lorenz, H., Der Baugrund als Federung und Dämpfung schwingender Körper, Bauingenieur, p. 11, 1950.

Cold Cathode Gas Tubes for Telephone Switching Systems

By M. A. TOWNSEND

(Manuscript received September 4, 1956)

Cold cathode gas tubes perform both switching and memory functions in telephone switching systems. One measure of the performance of a switching diode is the switching voltage gain, defined in terms of the characteristics of the device. Some of this gain must be sacrificed in order to increase the switching speed in a way which is analogous to the gain-bandwidth property of a conventional amplifier. In this paper, methods of achieving a high switching-voltage gain are described in terms of the gas discharge processes. An example is given of an application of these principles to a tube for use as a switch in series with the talking path in an electronic telephone switching system.

INTRODUCTION

Gas discharge tubes have found extensive use in telephone switching and other digital systems. Most of these applications take advantage of the fact that both switching and memory can be provided by a single gas discharge device. The switching characteristics result from the fact that the device is an essentially open circuit when the gas is not ionized and a closed circuit when the gas is ionized. The memory function is possible because the tube can be held in a high current condition, once it is ionized, by a voltage which is too low to initiate this conduction. Thus a triggering signal which ionizes the tube is "remembered" until the holding voltage is removed and the tube is allowed to de-ionize.

In some applications, gas tubes are used as switching devices in series with voice frequency circuits. For this purpose, the tube must offer a low impedance to audio frequency signals in addition to meeting requirements of switching and memory.

This paper first describes some switching characteristics of gas tubes considered as circuit elements. Desirable performance objectives are established in terms of these device characteristics. Following this, physical processes within the tube are described as they relate to circuit per-

formance. Finally, a description is given of a new talking path tube in which improved switching and transmission performance have been achieved.

EXTERNAL SWITCHING PROPERTIES

From the point of view of the external circuit, a gas diode may often be treated as a device which is an open circuit so long as the applied voltage is low, and which becomes a conductor if the applied voltage is raised above a threshold or "breakdown" value for a sufficient length of time. When the tube is conducting currents of the order of a few tens of microamperes or higher, the voltage is relatively independent of current and has a value, referred to as the "sustaining voltage", which is less than the breakdown voltage.

Although the details of actual circuits differ, it is possible to illustrate some important switching principles by the simplified circuit of Fig. 1(a). A gas diode is shown with the cathode connected to a bias voltage E through a load resistor. A signal source, assumed to have zero internal impedance is connected to the anode. The output voltage waveform corresponding to a pulse input is shown in Fig. 1(b). Note that after a time delay, t , the output rises to a voltage that differs from the total applied voltage by an amount equal to the sustaining voltage of the tube. The memory function is illustrated by the fact that the output signal remains after the input signal is removed.

It is often desired to use the output signal resulting from the triggering of one tube to switch one or more additional tubes. Since the input signals can be a few microamperes and the output signal can be tens of milliamperes, a large current gain is available from a gas tube. In many applica-

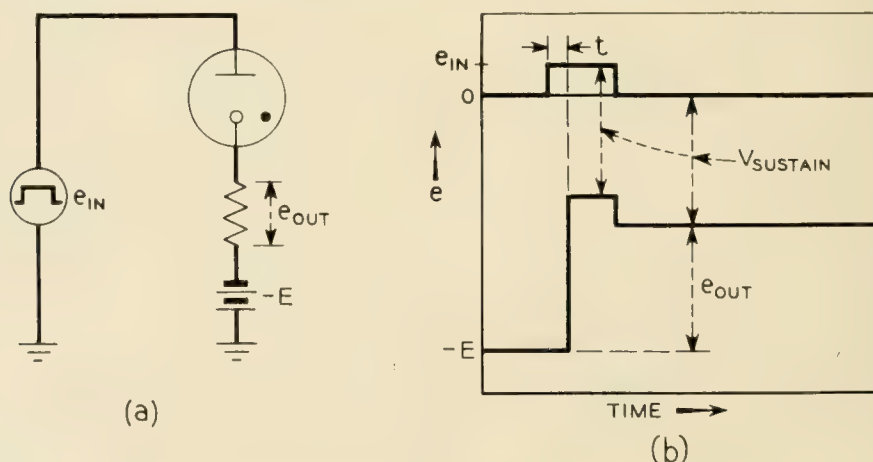


Fig. 1 — Simplified gas diode switching circuit.

tions, however, it is desired to apply the output voltage directly to other tubes without impedance transformation. In this case, voltage gain is of more interest than current gain. The maximum voltage gain per stage, defined as the maximum output voltage divided by the minimum input signal, is limited by variation in tube characteristics as will now be shown.

The bias voltage E of Fig. 1 is expected never to cause breakdown during times when the input voltage is zero. This establishes the upper limit of E as

$$| E | \leq V_{B \text{ min}} \tag{1}$$

where $V_{B \text{ min}}$ is the minimum breakdown voltage at any point in the life of any tube to be used in the circuit. As the bias voltage E approaches breakdown, the input signal voltage required for triggering approaches zero, and if there were no variation in breakdown voltage or bias voltage, and no noise voltages, the gain could be made to approach infinity.

The input signal, added to the bias, must be made large enough always to cause breakdown. Thus the minimum input signal is determined by $V_{B \text{ max}}$, the maximum breakdown voltage at any point in the life of any tube to be used in the circuit:

$$e_{\text{in}} \geq V_{B \text{ max}} - E \tag{2}$$

Combining (1) and (2)

$$e_{\text{in}} \geq V_{B \text{ max}} - V_{B \text{ min}}$$

or

$$e_{\text{in}} \geq \Delta V_B \tag{3}$$

where ΔV_B is the maximum variation in breakdown voltage among all tubes to be used in the circuit.

The output signal is the difference between the bias voltage and the sustaining voltage of the tube:

$$e_{\text{out}} = E - V_{\text{sus}} \tag{4}$$

The minimum output voltage corresponds to the maximum sustaining voltage, $V_{\text{sus max}}$. It is this value that must be used in calculating the maximum gain per stage as limited by the tube characteristics. The gain is then calculated as

$$G = \frac{e_{\text{out}}}{e_{\text{in}}} = \frac{E - V_{\text{sus max}}}{\Delta V_B} = \frac{V_{B \text{ min}} - V_{\text{sus max}}}{\Delta V_B} \tag{5}$$

This gain cannot be realized in practice because additional allowances

must be made for the variation in power supply voltages and protection against noise. In some cases the need for higher speed reduces the gain still further, as will now be shown.

In Fig. 1 a delay, t , is indicated between the application of the triggering signal and the appearance of the output signal. Part of the delay is statistical in nature and part is occasioned by the building up of ionization within the tube. As will be discussed later, the delay can be reduced by tube design techniques. However, for any given tube, the delay is a function of the excess of the triggering voltage over the breakdown voltage. The larger this overvoltage, V_{ov} , the shorter is the breakdown delay. Since this overvoltage must be added directly to the input signal, the gain is reduced.

Although not shown in Fig. 1, the tube is turned off by applying a signal that reduces the anode-to-cathode voltage below the sustaining value. This turn-off signal must have sufficient duration so that the tube does not again break down at the return to normal bias conditions. If the turn-off pulse duration is less than that needed for complete recovery, the effective breakdown voltage is reduced. Equation (5) can be modified to show the effect of this reduction in turn-off time by defining a quantity V_r , the reduction in breakdown voltage resulting from incomplete recovery of the tube. The combined effects of V_r and V_{ov} are then

$$G = \frac{(V_{B \text{ min}} - V_r) - V_{\text{sus max}}}{\Delta V_B + V_{ov}} \quad (6)$$

Equation (6) shows that faster turn-on obtained by increasing the over voltage V_{ov} and faster turn-off obtained by allowing for decrease in breakdown voltage by an amount V_r , both result in a reduction in voltage gain. Thus the familiar trade of speed for gain extends to gas tube switching circuits. Summarizing, it can be seen by (5) that constant breakdown voltage and large difference between breakdown and sustain are desirable switching properties. Also, as shown in (6), the tube should be designed so that the overvoltage needed to cause fast breakdown is small and the recovery of breakdown voltage after the tube is turned off is fast. It is useful to consider now the internal physical processes of a cold-cathode glow-discharge tube in order to see how the desired external properties can be obtained.

PHYSICAL PROCESSES OF A COLD CATHODE GLOW DISCHARGE

Since the gas particles are neutral and the cathode does not spontaneously emit electrons, current flow requires an auxiliary supply of charged particles. A small amount of radioactive material to ionize some

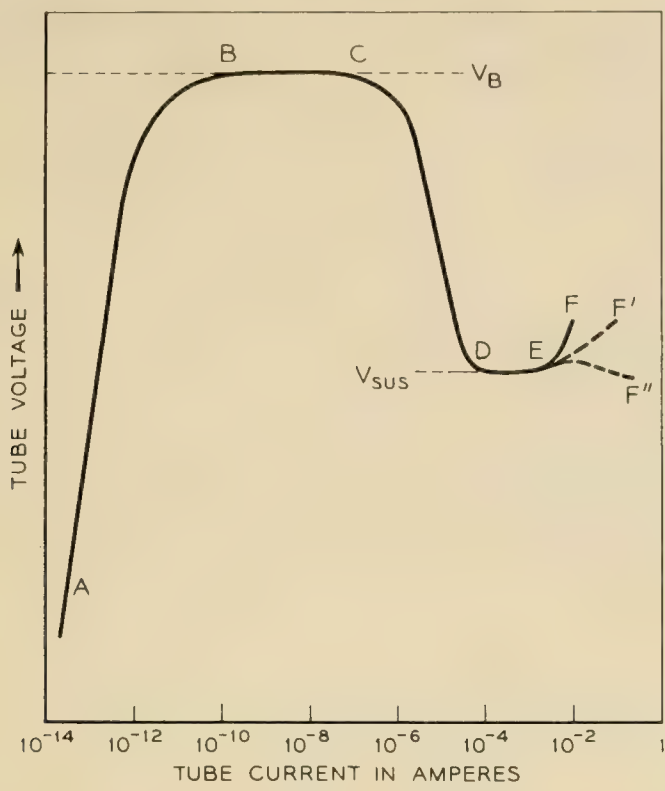


Fig. 2 — Voltage versus current curve of typical gas diodes.

of the gas or very small photoelectric emission of electrons from the cathode are commonly used for this auxiliary supply. A typical voltage-current curve is shown in Fig. 2. At low voltage, the current is very small, often being in the range of 10^{-14} ampere or less. The current increases with the voltage because collisions of electrons with neutral gas atoms produce additional excitation and ionization in the gas. Some of the new ions and excited gas atoms release new electrons by secondary emission when they strike the cathode.

The rate of increase of current with voltage depends on the kind and the pressure of the gas filling, the cathode material, and the tube geometry. An important characteristic of the gas is defined by an ionization coefficient η , which represents the number of new electrons (and ions) produced by a single electron moving through the gas a distance corresponding to one volt of potential difference.¹ This coefficient is a function of the kind of gas and of the quantity E/p_0 where E is the voltage gradient and p_0 is the normalized gas pressure. The fact that there is an optimum E/p_0 at which η is a maximum will be important to later discussion. The electron current at the anode, i_a , produced by gas amplification of a photoelectric current i_0 at the cathode is¹

$$i_a = i_0 e^{\int_{V_0}^V \eta dV}$$

(7)

where V is the anode voltage and V_0 is the initial voltage through which the electrons must travel before they can ionize.

The ions produced in the space by this process flow back toward the cathode. The ion current i_p resulting from this process is

$$i_p = i_0(e^{\int_{V_0}^V \eta dV} - 1) \quad (8)$$

As mentioned above, new electrons are released at the cathode by positive ions, neutral atoms excited to a metastable state, and photons generated in the gas. These secondary processes can be grouped together by defining a coefficient γ as the number of new electrons released at the cathode by all of these processes for each positive ion generated in the cathode-anode space. Thus each electron passing from cathode to anode, on the average, results in the release of M new electrons where

$$M = \gamma(e^{\int_{V_0}^V \eta dV} - 1) \quad (9)$$

Each new electron from the cathode is also amplified in the gas so that after n multiplication cycles, the electron current at the anode is²

$$i_n = i_0 e^{\int_{V_0}^V \eta dV} (1 + M + M^2 + \cdots + M^n) \quad (10)$$

When M is less than unity and $n \rightarrow \infty$ (the equilibrium state), (10) reduces to a steady state value of

$$i = \frac{i_0 e^{\int_{V_0}^V \eta dV}}{1 - M} \quad (11)$$

Since the current is dependent on the initial current i_0 , the discharge is said to be non-self-sustaining. This corresponds to the portion AB of the curve of Fig. 2.

If the applied voltage is made high enough, the multiplication factor approaches unity, the current of (11) becomes independent of the initial current, and the tube is said to have broken down. This condition corresponds to the horizontal portion BC of Fig. 2. To control this breakdown voltage, the cathode secondary emission coefficient γ and the gas ionization coefficient η must be controlled.

The secondary emission coefficient is highly sensitive to the surface conditions of the cathode. Pure metals such as molybdenum are often preferred to coated surfaces because they permit highly stable and reproducible emission. With the cathode surface determined, the breakdown voltage can be adjusted by changing the gas filling and tube geometry. Fig. 3 shows the breakdown voltage for a tube having parallel-plane

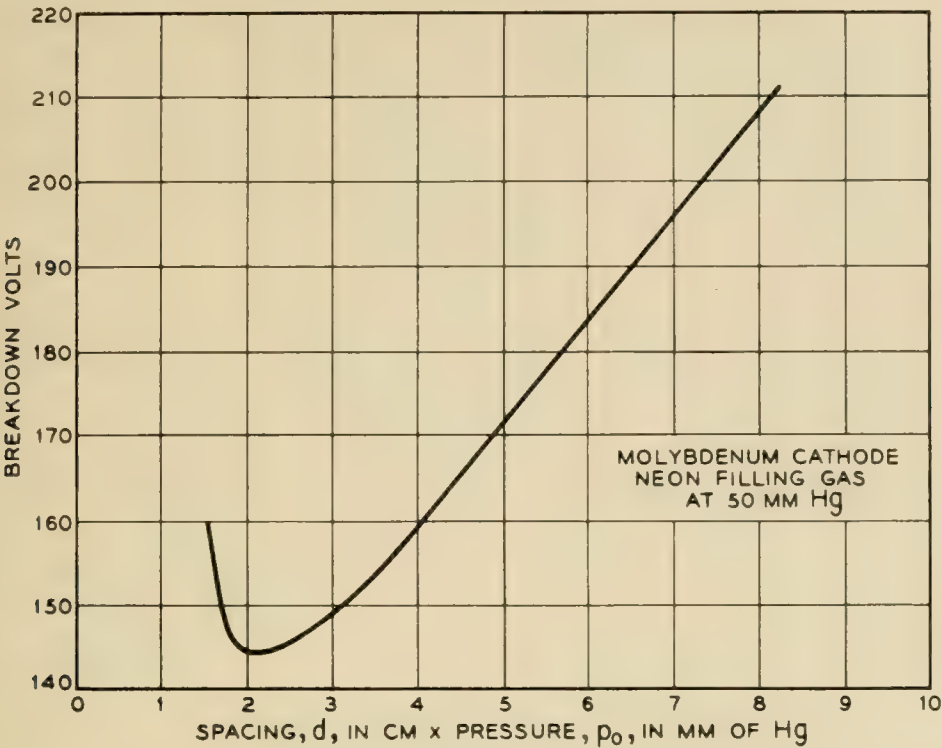


Fig. 3 — Breakdown voltage as a function of spacing and pressure for parallel plane anode and cathode.

anode and cathode geometry, a molybdenum cathode, and neon filling gas. The curve is plotted as a function of the product of pressure p_0 in mm Hg and electrode separation d in cm. Approximately the same plot would obtain for other pressures because both η and γ are functions of (E/p_0) and, for uniform fields, E is simply the voltage divided by the separation

$$\frac{(E)}{(P_0)_{\text{at breakdown}}} = \frac{V_{Bd}}{d} \frac{1}{p_0}$$

(12)

Since the variation of γ with E/p_0 is small and may be ignored in this elementary discussion, the minimum breakdown voltage corresponds very nearly to the optimum value of the ionization coefficient η . At spacings or pressures less than optimum, η is reduced because some electrons strike the anode without colliding with gas atoms. At spacings or pressures greater than optimum, η is reduced because electrons do not gain enough energy between collisions to ionize efficiently.

It can be seen that a way of meeting the switching requirement of constant breakdown voltage would be to design the tube to operate at the minimum of Fig. 3. Minor changes in spacing or filling pressure from one tube to another and changes in pressure with tube operation would result in small changes in breakdown voltage. The advantages of op-

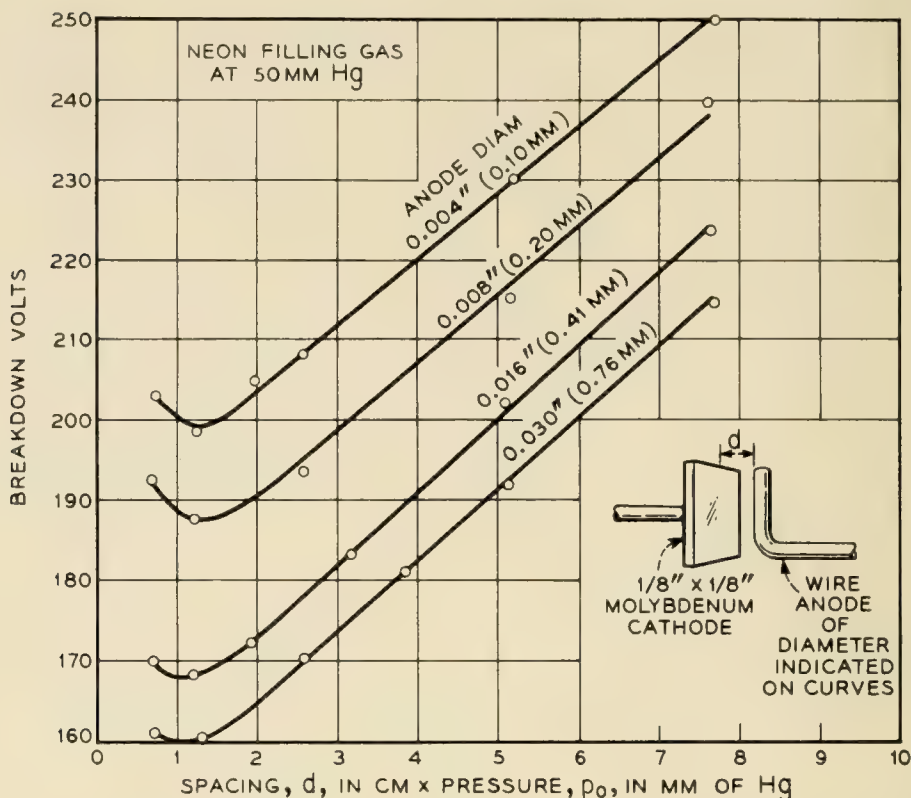


Fig. 4 — Breakdown voltage as a function of spacing and pressure for small wire anodes parallel to cathode surface.

eration at the minimum of the p_0d curve can be retained and the breakdown voltage made higher by resorting to non-uniform geometry. Typical curves are shown in Fig. 4. The cathode was a small rectangular plate and the anode was a wire placed parallel to the cathode surface. It is seen that as the anode diameter is decreased, the minimum of the breakdown curve is increased. The practical limit is set by mechanical stability of the anode and by transmission requirements, as will be discussed later.

The rise in minimum breakdown voltage as the anode size is reduced can be explained on the basis of the distortion of the electric field. Near the cathode, E is lowered, and near the anode, E is increased, as compared to the parallel plane case. If the spacing is adjusted for optimum E/p_0 with parallel planes, then η is necessarily less than optimum for the distorted fields.

Returning now to Fig. 2, we note that, as the current is increased beyond breakdown, the tube voltage falls to a lower sustaining value and again is relatively constant with current. This lower voltage corresponds to the development of a space-charge layer of positive ions near the cathode and an increased voltage gradient at the cathode. This higher

field results in an increase in the ionization coefficient η , and, in some cases,³ a larger effective value of the secondary emission coefficient γ . This is because electrons released by the secondary emission processes may strike neutral gas atoms and be reflected back to the cathode. A higher gradient increases the probability of escape of such an electron. Thus the multiplication factor M of (9) can equal unity at a lower total applied voltage.

Practical tubes filled with neon or argon gas have sustaining voltages near 100 volts when pure molybdenum or tungsten cathodes are used. Cathodes coated with barium and strontium oxide may sustain at 60 volts. However, since this lower sustain is accompanied by a lower breakdown voltage, the difference between them is not increased. Also, since the coated cathode surface is more variable between tubes and with tube operation, the switching voltage gain may be reduced with such cathodes.

The gas pressure and cathode geometry determine the length of the flat portion DE of Fig. 2. Over this current range, the area covered by the glow discharge increases with current until at E the cathode is completely covered. Increasing the cathode area or gas pressure increases the total current required for coverage. At still larger currents, the sustaining voltage increases rapidly as indicated by the solid curve EF . Broken curve EF' applies to a special cathode geometry called a hollow cathode.⁴ Such a cathode may be formed by the interior of a cylinder or by placing two plane cathodes close together so that the negative glow regions overlap. Under this condition electrons, ions, and excited atoms generated near one cathode can aid in current flow from the other cathode. Dotted curve EF'' applies to a particular form of hollow cathode⁵ in which cathode shape and gas pressure have been selected to give a negative slope in the high current region. This negative slope represents a negative resistance and permits audio-frequency signals to be transmitted through the tube without loss.

Anode effects have not been discussed. In general, the anode shape and location do not affect the sustaining voltage or the ability of the discharge to transmit audio frequency unless the anode-cathode spacing is too large. The basic requirement is that the anode should be large enough to intercept enough electrons to carry whatever current is required by the external circuit. Even a small anode placed near the cathode space-charge region can meet this requirement. Thus the sustaining voltage of a tube designed to have a breakdown voltage near the minimum of Fig. 3 or Fig. 4 will not in general be sensitive to the anode size or shape.

The transition from low current to high current in a gas diode can thus be thought of as the process of introducing a space charge of positive ions in the region near the cathode. This is done by raising the voltage temporarily above the breakdown value. To switch back to the low current, it is necessary to decrease the multiplication factor M below unity by temporarily lowering the voltage and allowing the ions and excited atoms to diffuse out of the cathode and anode region. Both the turn-on and the turn-off processes impose time restrictions on the switching characteristics.

The multiplication factor M of (9) applies to an average process. Thus, even though M is greater than unity, it is possible that the ionization and excitation produced in the gas by any individual electron may not release a new electron at the cathode. It is therefore necessary on the average to wait for more than the time between initiating electrons before the discharge starts to build up.

The average statistical delay is then equal to the average time between successful starting events. If N_0 photoelectrons per second are emitted from the cathode and W is the fraction of these which successfully initiate a discharge, the average statistical delay is⁶

$$t_{AV} = \frac{1}{WN_0} \quad (13)$$

The fraction W would be expected to increase with an increase in the multiplication factor M and hence with the overvoltage above breakdown. It has been shown theoretically and experimentally⁷ that this is the case. For voltages only slightly in excess of breakdown, i.e., small overvoltages, V_{ov} , the expression for average statistical delay can be approximated by

$$t_{AV} \approx \frac{k_s}{V_{ov}} \quad (14)$$

In practical tubes with overvoltages of 10 volts, the average statistical delay may be of the order of milliseconds with radioactive sources of ionization. Short delays of the order of microseconds are obtained by providing an auxiliary "keep-alive" discharge to a separate electrode or by illumination that provides a photoelectric current in the range of 10^{-12} amperes.

A formative delay in breakdown also occurs because time is required for current to build up to the final value. This time is equal to the product of the number of multiplication cycles and the time per cycle. The number of multiplication cycles required is reduced as multiplica-

tion factor M is increased with increasing overvoltage. The multiplication factor M includes electrons released at the cathode by slow moving metastable gas atoms as well as those released by the faster positive ions. At very low overvoltages, these slow components must be included before the current can build up.⁸ At higher overvoltages enough positive ions are produced so that M is greater than unity without waiting for the slow components. Thus the effective time per multiplication cycle is reduced with increasing overvoltage. Since the number of cycles and the time per cycle are both decreased the formative delay decreases rapidly with increasing overvoltage. A typical formative delay for a neon filled, molybdenum cathode switching tube at 5 volts overvoltage might be of the order of 100 microseconds.

APPLICATION TO A TALKING-PATH SWITCHING DIODE

The principles discussed above have been applied in the development of a cold-cathode gas diode for use as a switch in series with the speech path in an electronic switching system. The objectives were a switching voltage gain as high as possible, a breakdown time of less than a few hundred microseconds, and a low transmission impedance for audio-frequency signals.

A sketch of one version of the resulting tube is shown in Fig. 5. The cathode is a molybdenum rod which has a small hollow cathode portion in the upper end. The anode is a small molybdenum wire placed near the minimum breakdown distance and slightly to one side of the opening in the end of the cathode. A barium getter is flashed to one side of the bulb wall and a small tungsten wire spring is arranged to make electrical contact with the getter flash. A neon filling gas at a pressure near 100 mm Hg is used.

The cathode geometry has several interesting properties. It was found that the shape of a cylindrical hollow cathode is unstable at very high current densities and that it will rapidly grow into a spherical cavity with a small orifice.* Typical dimensions are a sphere diameter of 0.030 inch and an orifice diameter of 0.008 inch. At an operating current of 10 milliamperes, the current density in the orifice is of the order of 50 amp/cm². Once the sphere has stabilized it will operate many thousands of hours with relatively small changes in shape. The transmission properties of the stabilized spherical cavity cathode are similar to the earlier negative resistance hollow cathode tubes.⁵ Typical im-

* This cathode was developed by A. D. White of Bell Telephone Laboratories and will be described more completely by him in a forthcoming publication.

pedance values are 300 ohms negative resistance and 50 ohms inductive reactance at 10 milliamperes operating current, with a superimposed audio-frequency signal of 3,000 cycles per second.

Even though the cathode geometry is stable, some cathode material escapes through the orifice and will rapidly collect on an anode placed directly over the opening. It is therefore necessary to locate the anode to one side of the orifice. The extremely high ionization density near the cathode orifice allows considerable flexibility in anode location without affecting the sustaining voltage or destroying the negative resistance.

High switching-voltage gain is obtained by using a small anode formed by a 0.005-inch diameter molybdenum wire placed perpendicular to the end of the cathode at a spacing of approximately 0.005 inch. Breakdown voltage is nominally 190 volts with a range of ± 10 volts over all tubes and over the nominal operating life of 4,000 hours. The sustaining volt-

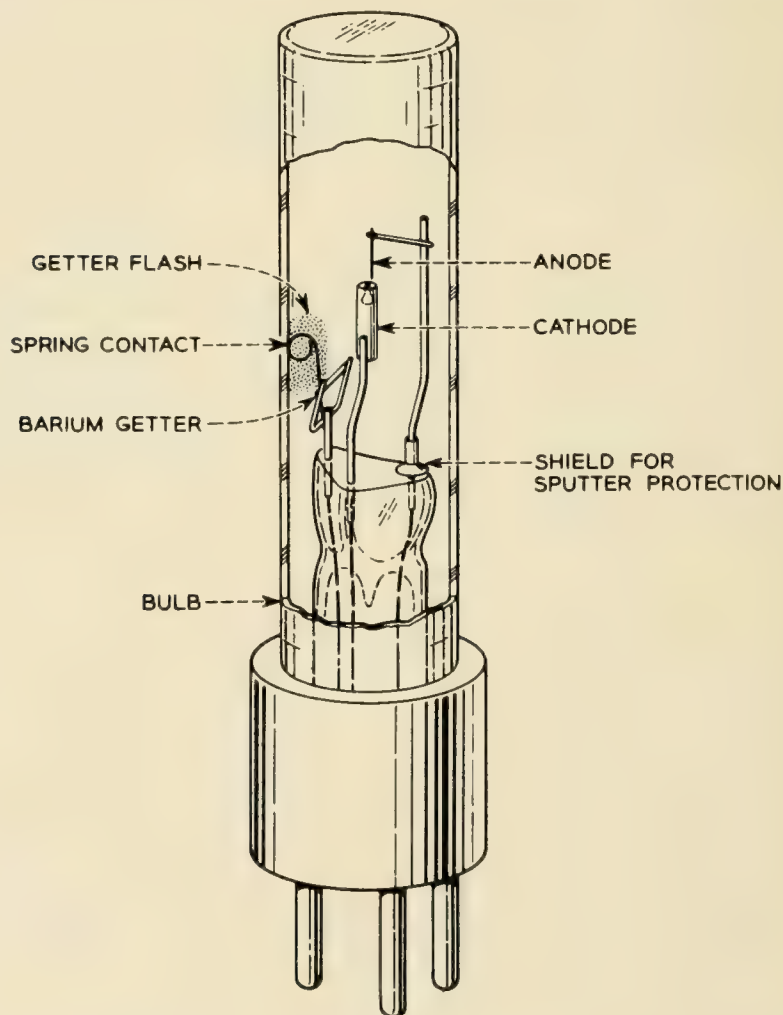


Fig. 5 — A talking-path switching diode.

age at the operating current of 10 ma is 99 ± 2 volts. Thus the switching gain from (5) is $(180-101)/20$ or 3.9. In practice, switching is often done without allowing the full 100 milliamperes operating current to flow. Under these conditions, the sustaining voltage may be 10 or 15 volts higher, with a consequent reduction in switching voltage gain.

Short breakdown times were desired for this tube. It was not desirable to use enough radium to obtain the needed initial ionization, since it is expected that large numbers of these tubes will be concentrated in a relatively small space. Also, the molybdenum cathode does not emit photoelectrons unless short wavelength ultraviolet illumination is used. The solution chosen was to use the barium getter flash as an auxiliary photocathode. An electrical contact is made to the getter deposit and this is connected through a high resistance to the main cathode. Visible light or long wave ultraviolet light is readily transmitted through the bulb and produces photoelectric current in the auxiliary gap. This current is amplified by the gas, but remains a non-self-sustaining discharge. Currents of 10^{-10} amperes are readily available with a few foot-candles of illumination. This current is too small to affect the breakdown voltage of the main gap, but produces enough residual ionization to allow breakdown times of the order of 100 microseconds to be obtained with a few volts overvoltage. The high resistance connection to the main cathode may be of the order of 20 to 50 megohms. It protects the photocathode from deterioration which might result from high currents when the main gap is conducting.

Recovery of breakdown voltage following conduction is rapid. Measurements indicate that the breakdown voltage is within the limits of 190 ± 10 volts in less than 500 microseconds. The relatively high gas pressure and close spacings speed up the deionization process.

The tube described has not been designed for large scale manufacture although several hundred models have been made and tested to establish the feasibility of the design.

SUMMARY

Some useful switching properties of gas diodes can be described by defining the switching-voltage gain. This gain is shown to be equal to the difference between the breakdown and the sustaining voltage divided by the variation in the breakdown voltage. The gain is reduced if faster switching times are required.

The switching-voltage gain is discussed in terms of the physical processes in a gas discharge. It is shown that a high gain can be obtained by using an inefficient anode operating at the minimum of the curve

of breakdown voltage versus the product of gas pressure and anode distance.

A tube is described which uses these principles to achieve a high gain over a useful operating life of 4,000 hours, and which has a negative resistance to audio frequency signals superimposed on the dc operating current. Fast switching is obtained by an auxiliary photoelectric cathode formed by making an electrical connection to a barium getter flash. Satisfactory tube operation has been obtained for continuous operation for times which are equivalent to 20 to 40 years of intermittent operation in switching systems.

ACKNOWLEDGEMENT

The author is indebted to many members of the gas tube and switching systems development groups at Bell Telephone Laboratories. Among these special mention should be made of A. D. White who originated the cavity hollow cathode and V. L. Holdaway, B. T. McClure, A. M. Wittenberg and C. Depew who made important contributions to the successful development of the tubes.

BIBLIOGRAPHY

1. M. J. Druyvesteyn and F. M. Penning, *Rev. Mod. Phys.*, **12**, p. 97-102, 1940.
2. *Ibid.*, page 105.
3. R. N. Varney, *Phys. Rev.* **93**, p. 1156, 1954.
4. Reference 1, p. 139.
5. M. A. Townsend, W. A. Depp, *B.S.T.J.*, **32**, pp. 1371-91, 1953.
6. Reference 1, p. 116.
7. F. G. Heymann, *Proc. Phys. Soc.*, **63**, Sec. B, 1950.
8. H. L. Von Gugelberg, *Helvetica Physica Acta*, **20**, pp. 307-340, 1947.

Activation of Electrical Contacts by Organic Vapors

By L. H. GERMER and J. L. SMITH

Unreproducibility of earlier work on the erosion of relay contacts has been traced to the effects of organic vapors in the atmosphere. Carbon from decomposition of these vapors greatly alters the conditions under which an electric arc can be initiated and can be sustained. The importance from the standpoint of erosion comes from the fact that for many circuit conditions contacts activated by this carbon cannot be protected against severe arcing by any conventional capacitance-resistance network. This paper reports investigations which have enabled us to understand the activation of contacts by organic vapors.

TABLE OF CONTENTS

Introduction	770
<i>Part I. Electrical Effects</i>	772
1. Observations on Activation.....	772
1.1 Striking Field.....	774
1.2 Arc Voltage.....	775
1.3 Minimum Arc Current.....	777
1.4 Erosion.....	778
1.4(a) Palladium and Platinum.....	778
1.4(b) Silver and Gold.....	779
2. Interpretation of Activation.....	780
2.1 Striking Field.....	781
2.2 Arc Voltage.....	784
2.3 Minimum Arc Current.....	785
2.4 Erosion.....	786
2.4(a) Palladium and Platinum.....	786
2.4(b) Silver and Gold.....	788
2.4(c) Anode Arcs and Cathode Arcs.....	790
3. Recapitulation.....	794
<i>Part II. Activating Carbon</i>	796
4. Composition of Activating Powder and Rate of Production.....	796
4.1 Composition.....	796
4.2 Rate of Production.....	797
5. Surface Adsorption.....	798
5.1 Benzene Molecules on Contact Surfaces.....	798
5.2 Inhibiting Surface Films.....	801
5.3 Alloys.....	802

6. Activation in Air.....	803
6.1 Burning of Carbon.....	804
6.2 Diffusion of Activating Vapor.....	806
6.3 Sputtering and Burning in a Glow Discharge.....	807
6.4 "Hysteresis" Effects.....	809
7. Brown Deposit.....	810
7.1 Composition.....	810
7.2 Rate of Production.....	811
7.3 Brown Deposit and the Carbon of Activation.....	811

INTRODUCTION

Contamination of surfaces by organic vapors is a subtle factor that influences the electrical erosion of relay contacts. Because of this contamination, contacts in the telephone plant sometimes erode very much more than one would expect from simple laboratory life tests. This caused considerable confusion until about 1945 when the influence of organic vapors was recognized. The term "activation" is used here to describe changes in the surfaces of electrical contacts which give rise to greater arcing when an electrical circuit is completed or broken than would occur if the metal surfaces were clean.* Although its cause is generally carbon from organic vapors, there are occasionally other causes. This paper is an account of recent research¹ on activation produced by organic vapors.†

It has been found that the carbon that causes activation is formed on the electrode surfaces by decomposition of adsorbed organic molecules. Microscopic examination of contacts gives a very sensitive way of detecting incipient activation, since the carbon can easily be seen before any electrical effects are observed. The minimum amount of carbon necessary for activation is of the order of 0.05 microgram.

Activation has been produced on noble metals only, and only by unsaturated ring compounds. When experiments are carried out on clean noble metal surfaces under controlled conditions which do not permit burning of carbon, it is found that the amount of carbon formed by an arc corresponds to approximately a monolayer of organic molecules on the area heated by the arc. After a surface has become active, the amount of carbon formed by each arc is considerably increased and corresponds to the decomposition of several monolayers of molecules. In air, the situ-

* The term "activation" has sometimes been used heretofore to signify enhanced erosion resulting from organic vapors. This is a different definition from that used in this paper, due to the fact that in some cases, long sustained arcs produce less erosion than arcs of shorter duration. This is often true for silver surfaces, as described below. In a case of this sort, a surface may have a great deal of carbon on it and be very "active" by our definition, when it would be considered not active at all by the definition that relates activation to rate of erosion.

† Other causes of activation will not be considered here. See Reference 2, page 961.

ation is much complicated by burning of carbon and by the impedance offered by air to diffusion of molecules to the electrode surfaces. Because of these complicating factors, activation will not occur in air if the vapor pressure is too low or if the time between arcs is too short.

Arcs at the making and breaking of clean contacts — clean in the sense that they are free from carbon — produce transfer of metal from one contact to the other with a resulting pit and mound of about equal volumes. The situation is greatly changed by carbon. The presence of carbon causes increased arcing, alters the characteristics of the arcs, and greatly changes the resulting erosion both in character and amount. With carbon present, some or even all of the eroded metal does not stick to the electrodes, and there is often loss of metal from both of them, the missing metal turning up mixed with carbon in a loose black powder. With carbon on the surfaces, successive arcs occur at different places, and the resulting erosion tends to be smooth with the electrodes worn down uniformly all over their surfaces. This is because each arc burns off carbon at its center, while it produces more around its periphery where the metal is cooler, and each new arc strikes on a newly carbonized surface.

Every arc, of either the active sort or of the “inactive” sort which occurs at clean surfaces, is predominantly an arc *in metal vapor*. The active arcs, as well as the arcs at clean surfaces, are of one or the other of two quite distinct types which have been called, respectively, “anode arcs” and “cathode arcs” (Reference 3 and 4 which are concerned with palladium electrodes only). In an anode arc, most of the metal of the arc is vaporized from the anode by electron bombardment, but in a cathode arc the metal is supplied from the cathode by the explosion of small areas due to Joule heating by field emission currents of enormous densities flowing through them. In an anode arc, the erosion is predominantly from the anode, and in a cathode arc from the cathode.

Whether a particular arc is of the anode or of the cathode type is determined by the electrode separation and the contact metal. For palladium electrodes, an arc is an anode arc if the separation is less than about 0.5×10^{-4} cm, but a cathode arc if the separation is greater than this value. The corresponding critical distance for silver is 3 or 4×10^{-4} cm. The carbon particles producing activation permit breakdown at separations for which it would not occur in the absence of carbon, and thus favor cathode arcs. The critical distance of palladium is so small that all active palladium arcs are cathode arcs, with the greater loss of metal from the cathode. For silver, on the other hand, the critical distance is so large that active arcs at silver surfaces are in many cases anode arcs, with the greater loss of metal from the anode as in the case of inactive silver arcs.

PART I — ELECTRICAL EFFECTS

1. OBSERVATIONS ON ACTIVATION

Just how different the behavior of contacts can be in the clean or inactive condition, and in the active condition, is shown strikingly by the oscilloscope traces of Fig. 1. Each trace represents a plot against time of the voltage across a pair of protected relay contacts when the contacts are pulled apart to break a current through an inductive load. The circuits used with the two pairs of palladium contacts were identical, the only difference of any sort being in the conditions of the surfaces of the contacts produced by exposure of the second pair of contacts to organic vapor. In the trace of Fig. 1(a), the potential across the contacts

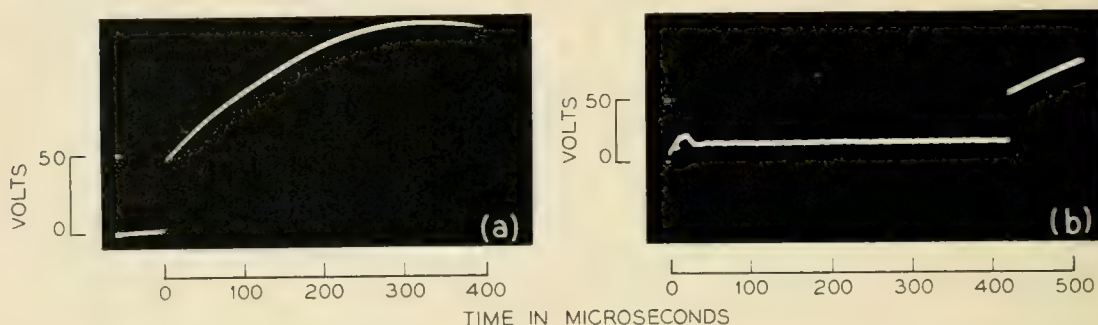


Fig. 1 — Oscilloscope traces of the voltage across relay contacts breaking a current of half an ampere through an inductive relay load. In each case a standard protective network of a $0.5\ \mu\text{f}$ capacitor in series with 100 ohms is in parallel with the contacts. (a) Clean or "inactive" contacts, with no observable arc. (b) Active contacts, with a sustained arc lasting 400 microseconds.

risers abruptly on break from zero to 50 volts, and then continues to increase as the capacitor of the protective network is gradually charged up; there is no arc or other discharge at the contacts. In the trace of Fig. 1(b), the potential rises on break to about 14 volts and remains at this value for 400 microseconds. This represents an electric arc that occurred across the contacts, the 14 volts being the potential characteristic of arcs over short distances at palladium electrodes (Reference 2, Table II). The energy dissipated at the contacts by this arc was about 25,000 ergs.

A number of worth while experiments upon activation can be carried out with no better method of measuring, or detecting, activation than the observation of oscilloscope traces like those of Fig. 1, or corresponding traces obtained from contacts discharging a small capacitor on closure (Reference 2, Fig. 1). One can find what organic vapors produce activation, and what metals can be activated. This can be extended to discover

what vapor pressure is needed and how the minimum pressure depends upon the conditions of the test; for example, upon the electrical test circuit. Some results of simple tests of this sort will be given before going on to present anything more fundamental about what is happening when contacts become active.

Activation is produced, in general, by repeated operation of a pair of contacts closing or breaking an electric circuit in air containing an organic vapor. At the beginning of a test of this sort, oscilloscope traces look the same as they would look in the absence of the vapor, but with continued operation, arcs may begin to occur when there were initially no arcs, or the durations of arcs may become greater. Activation is not produced by simply allowing contacts to stand idle in an activating vapor even for very long times, and even when the partial pressure of the vapor is extremely high. Nor is activation produced unless currents are made or broken. Operating contacts "dry" in a high pressure of an activating vapor does indeed make them temporarily active,* but this condition is unimportant from the standpoint of erosion. Under conditions that are effective in producing activation, the number of operations required before increased arcing can be detected is usually greater than 100, and often greater than 10^4 .

In general, noble metals can be activated by organic vapors and base metals cannot. Vapors of unsaturated ring compounds produce activation and other organic vapors do not. Tests have been made upon the metals Ag, Au, Cu, Pd and Pt which can be activated and upon Co, Cr, Fe, Mn, Mo, Ni, Sn, Ta, Ti and W which have not been activated.† The vapors of nearly 50 organic compounds have been tested, about half of them unsaturated ring compounds which produce activation, and about half other compounds which do not. (These are listed, in part only, in Reference 2, Table 1). A very large number of tests have been carried out upon benzene, limonene and styrene. For these three compounds the minimum partial pressures which just produce activation of silver electrodes in air under certain standard conditions, and of platinum electrodes in air, for which the results are the same as for silver, were found to be respectively 0.1, 0.03 and 0.003 mm Hg.

Some insight into activation is obtained by direct examination of active contacts. Black soot can always be seen on active contacts, and if they have been operating for some time in activating vapor, the amount of soot may be great enough to produce a visible deposit under-

* See the Section 7.3 on "Brown Deposit."

† *Ni* and *W* have been activated in the presence of an organic vapor in a container in which the pressure of air was reduced to 0.01 mm Hg, Reference 5, page 1090.

neath the electrodes. It is clear that the increased arcing of activation is caused by solid carbonaceous material made by decomposition of organic vapor and not by the vapor itself. Clean metal contacts can, in fact, be made to show all of the symptoms of activation by allowing soot from a flame to settle on their surfaces.*† Activation produced in this way is, of course, temporary, lasting only until the deposited soot is burned away.

When one looks for characteristics of arcs between active surfaces to which numerical values can be attached, four features come at once to mind — the electric field at which an arc strikes, the voltage across the arc after it is established, the minimum arc current (which is just the current at which the arc goes out), and finally, after the arc is over, the amount of metal that was gained or lost by each of the electrodes during the arc. All of these quantities have been measured for active arcs as well as for arcs at clean surfaces, and a brief summary of the results of the measurements is given here.

1.1 *Striking Field*

To measure the electrode separation at which an arc strikes between closing electrodes, relay contacts were operated repeatedly, discharging on each closure a capacitor charged to a measured voltage. An arc at each closure was assured by using short leads between the capacitor and the contacts to keep the circuit inductance very low. The time from the initiation of the arc to the touching of the contacts was measured on an oscilloscope.‡

Fig. 2 illustrates the results of measurements made by F. E. Hawthorth⁷ upon palladium electrodes closing at 30 cm/sec to discharge a very small capacitor charged to 50 volts. Before the start of the experiment the electrodes had been cleaned by repeated arcing in air, and the first experimental point represents the closure of these clean electrodes. All of the other measurements were made in air containing a fairly high partial pressure of limonene vapor. Each point plotted on the curve rep-

* Unpublished work of P. P. Kisliuk.

† It is interesting to point out that a surface is not made active by rubbing petrolatum upon it, although activation will occur very quickly if an electric current is made or broken at such a greasy surface, so that some of the grease is decomposed.

‡ For inactive arcs, it is necessary to make a correction for the height of the mound of metal thrown up by the arc (Reference 6, page 1136). After the contacts become active, there is no appreciable mound thrown up (at palladium surfaces), and the electrode separation at the initiation of the arc is calculated at once from the closure time and the previously measured electrode velocity. The height of the mound produced before the contacts are active was minimized by using a capacitance of only 40×10^{-12} farad.

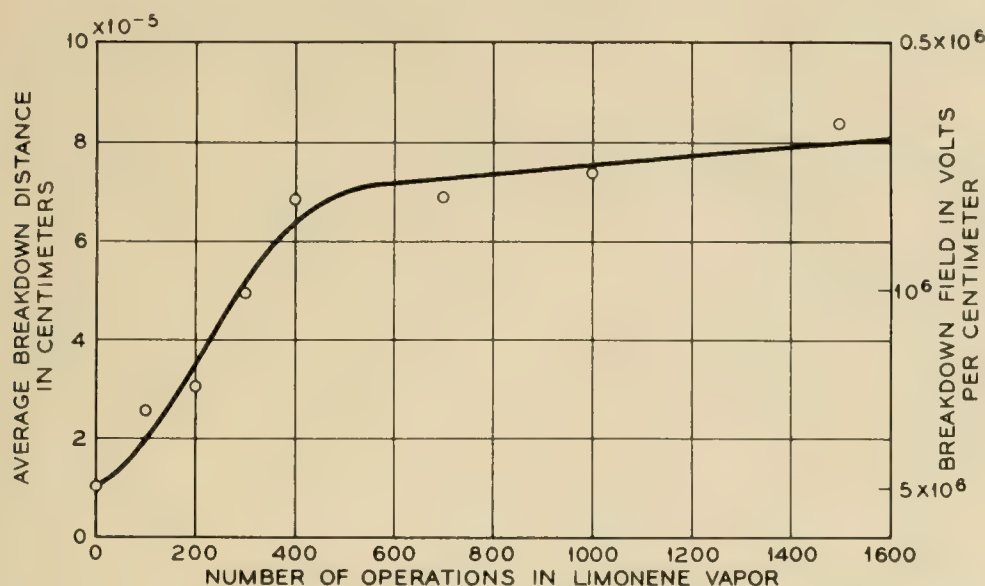


Fig. 2 — Breakdown distance, and apparent striking field, for arcing at relay contacts on closure in the presence of limonene vapor, plotted against number of operations. Each closure discharges a very small capacitor charged to 50 volts. Contacts are clean and inactive at the beginning of the test.

resents the average of 100 separate measurements. During tests of this sort, it was discovered that, with frequent microscopic examination of the electrodes, black sooty material could easily be seen after the first 30 closures, before certain evidence of activation could be obtained in any other known way. In Section 4.1, it is shown that this material is carbon.

The average electrode separation at which an arc struck between clean electrodes at the beginning of the curve of Fig. 2 was about 1×10^{-5} cm and after the electrodes became covered by sooty material, about 8×10^{-5} cm. The apparent striking field was decreased by activation from 5×10^6 to 0.6×10^6 volts/cm. When measurements were made at 250 volts, rather than 50 volts, the striking field in the active condition was only slightly higher, 0.8×10^6 volts/cm. Activation produces a lowering of the apparent striking field, regardless of the value of the applied voltage.*

1.2 Arc Voltage

The observed voltage across an arc at active palladium contacts agrees in general with that of palladium cathode arcs, which is about 16 (Reference 4, Fig. 7 and Reference 2, Table II), whereas the arc voltage of

* This apparent contradiction of the conclusion of F. E. Haworth⁷ is clarified in Section 2.1.

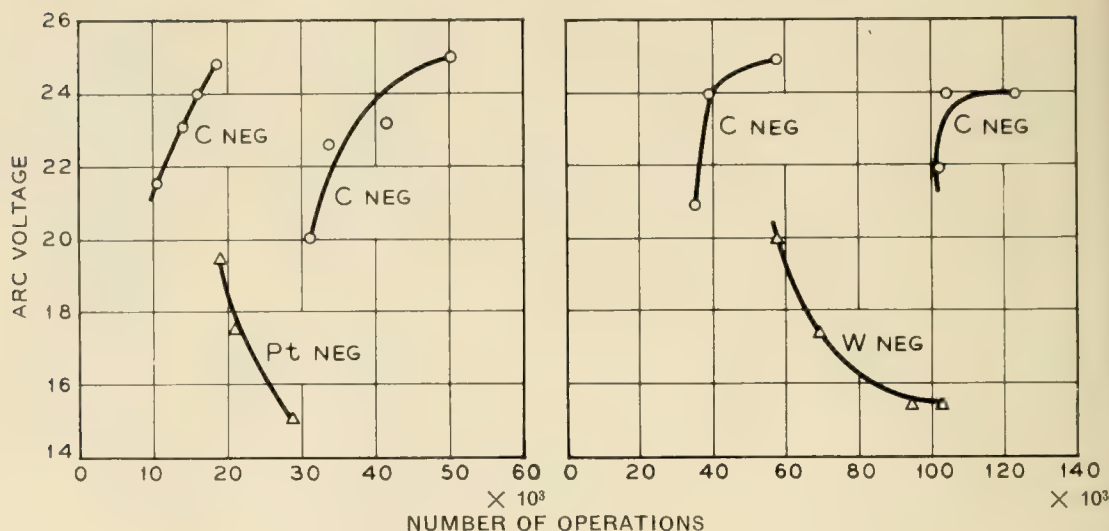


Fig. 3 — Measurements of arc voltage at cathode arcs between a carbon-platinum pair of electrodes, and between a carbon tungsten pair — at successive reversals of striking potential.

carbon is much higher and quite variable in the range from 20 to 30. One is tempted to conclude from this that the vapor in an arc between active palladium contacts is predominantly the metal of the electrodes, not carbon vapor. A more sound conclusion, however, as will be pointed out later, is that the source of electrons on the cathode of an active arc is palladium metal rather than carbon.*

When contact surfaces are very heavily carbonized, an arc voltage substantially higher than that characteristic of the metal of the electrodes is sometimes observed for a short time at the beginning and at the end of an active arc occurring at the discharge of a capacitor into an inductive circuit. An example of this is shown in the oscilloscope trace of Fig. 4. The higher arc voltage at the beginning of this arc, when the current was extremely small, is interpreted as the initiation of the arc between carbon surfaces, and the enhanced voltage at the end is evidence

* A very simple experiment has been carried out which proves conclusively that the character of an arc of the type which we call a cathode arc (see below, and Ref. 3 and 4) is determined by the properties of the cathode, and not by those of the anode. This is perhaps self evident, but a direct test is reassuring. The test is simply the observation that, for an arc of the cathode type between electrodes of different materials, the arc voltage is substantially the same as it would be if both electrodes were of the cathode material. The test is made by reversing the potential between the electrodes repeatedly, and after each reversal observing that the arc voltage changes gradually from that characteristic of the anode to that characteristic of the cathode. After each reversal the arc begins to clean from the cathode the anode material that was deposited there before the reversal, when what is now the cathode was the anode. Accompanying this cleaning, the arc voltage goes up or down until it reaches the value characteristic of the cathode itself. Measurements obtained in this way are reproduced in Fig. 3 for a carbon-platinum pair of electrodes and for a carbon-tungsten pair.

that when the current was again small the arc was localized at a new position on a fresh carbon surface so that it was again a carbon arc; during most of the arc time, when the current was larger, the cathode surface was maintained so free from carbon that the source of electrons at the cathode was palladium metal rather than carbon. For lightly carbonized surfaces, which may be just as active as judged by arc duration or any other test that we know of, no such enhanced arc voltage at the beginning or at the end of an arc has been observed. It may well be that for lightly carbonized surfaces the arc voltage is characteristic of carbon for a time too short to be detected by this crude means.

1.3 Minimum Arc Current

The current at which an arc goes out is readily found by observing on an oscilloscope the potential across closing contacts discharging a capacitor through a non-inductive resistor R . At extinction, the potential rises from the arc voltage v to that across the capacitor V_1 . The minimum arc current is then $(V_1 - v)/R$. An oscilloscope trace showing such a determination of minimum arc current at the arc initiation potential of 400 volts is reproduced as Fig. 5. (See also Reference 2, Fig. 5 and

Fig. 4 — Oscilloscope trace representing the voltage across an arc at the closure of very heavily carbonized electrodes. Discharge through an inductance of $10^{-4}h$ of a capacitor of $10^{-8}f$ charged to 50 volts. Near the beginning and near the end of the arc the source of electrons at the cathode was a carbon surface.

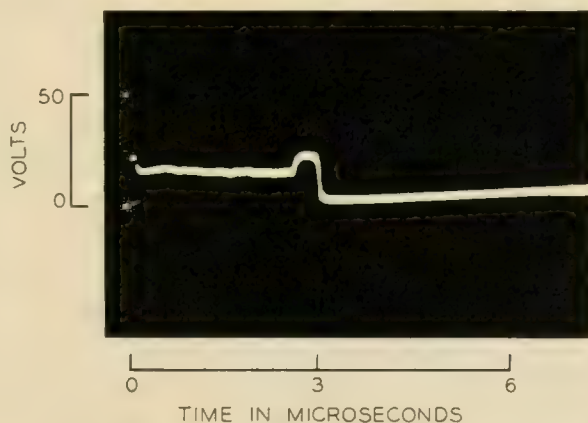
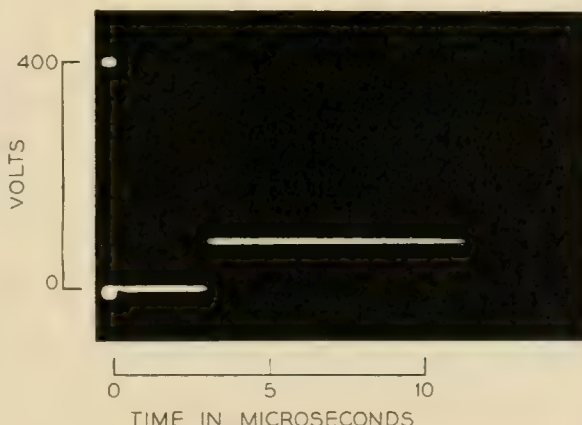


Fig. 5 — Voltage across clean palladium contacts when a capacitor charged to 400 volts is discharged through a resistor of 200 ohms. The closure arc went out at the minimum arc current 0.42 amp.



Reference 8, Fig. 1). The minimum arc current is much lower for active contacts than for inactive or clean contacts, and one can perhaps think of the decrease of the minimum arc current for noble metal contacts from a value of the order of 1 ampere for clean surfaces to 0.1 ampere or less for active surfaces as the chief characteristic of activation.

1.4 Erosion

Active contacts of palladium and of silver transfer metal in quite different ways. The transfer at active silver contacts is the more complex, and for this reason the transfer that occurs at active palladium contacts

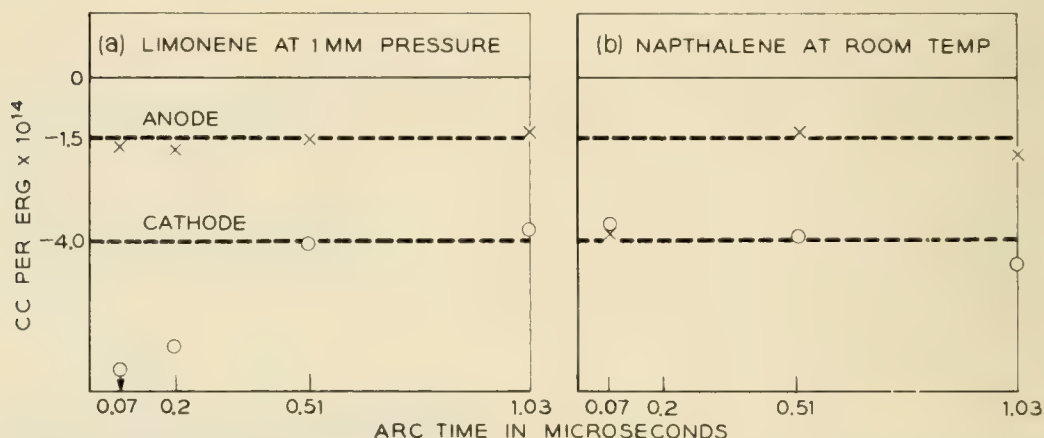


Fig. 6 — Results of measurement by weighing of the erosion of palladium electrodes produced by active arcs in limonene vapor (a) and in naphthalene vapor (b).

will be taken up first. The behavior of platinum is in general like that of palladium, and gold is like silver.

1.4 (a) Palladium and Platinum. It is found that arcing on closure at active contacts of palladium or platinum causes loss of metal at the cathode of the order of 4×10^{-14} cc/erg. The anode often loses metal also, but the loss at the anode is considerably less and may be zero in some cases. The results of two sets of measurements upon active palladium contacts are plotted in Fig. 6. These data represent changes in volume (calculated from weighings) per unit of arc energy after repeated arcs in limonene vapor at a vapor pressure of 1 mm Hg, Fig. 6(a), and in the vapor of naphthalene saturated at room temperature, Fig. 6(b). Tests were made by closing electrodes to discharge on each closure a properly terminated fixed length of cable charged always to 200 volts, to give in each case a constant arc current of 4 amperes, with the arc lasting for the time determined by the cable length. For the shortest arc time, the energy of

each of the individual arcs was 40 ergs, and for the longest arc time 600 ergs. The results indicate no significant variation of the erosion per unit of energy over this range.

There is some evidence that arcs at the break of active palladium surfaces give significantly lower cathode erosion per unit of energy (1 or 2×10^{-14} cc/erg) than do arcs at closure. The reason for the difference is not clearly understood, but widely different currents and electrode separations may be significant factors.

By examining contacts of palladium or platinum after many active arcs (on either break or closure), it is found that the erosion tends to be uniform over the surface, wearing each electrode down smoothly, with much less loss from the anode than from the cathode. This type of wear is quite different from that produced by arcs at clean surfaces. Erosion by arcs at clean surfaces always gives a mound of metal on one electrode, with a corresponding pit in the other; the loss of metal from one electrode is not appreciably greater than the gain by the other, the entire erosion consisting simply of transfer of metal between the contacts.

Now the inactive arcs at clean surfaces are known to be of two types which have been called "anode arcs" and "cathode arcs."^{3, 4} In anode arcs, the transfer of metal is about 4×10^{-14} cc/erg and is from anode to cathode, with a resulting pit in the anode and a matching mound on the cathode (Reference 9, page 1085–1086). In inactive cathode arcs, measurements made in the same way and not yet published have shown that the transfer is smaller — about 1×10^{-14} cc/erg — and is in the opposite direction, from cathode to anode, with a resulting pit in the cathode and a matching mound on the anode.¹⁰ It will be shown later, Section 2.4(a), that arcs at active palladium surfaces are of the cathode type, each individual arc being not readily distinguishable from an inactive cathode arc in the effect it produces on the cathode surface. The reason for the net cathode loss being greater in an active cathode arc than in a cathode arc at clean surfaces is due, at least in part, to some reverse transfer in an arc at clean surfaces.

1.4 (b) Silver and Gold. The erosion of silver surfaces is quite complex, and an adequate description of all of the phenomena encountered is reserved for later publication.¹⁰ A simplified description of the main features of the erosion of silver contacts is given here. Tests upon active gold contacts have been less extensive than upon active silver contacts, but as far as the observations go, gold has been found to behave just like silver.

At active silver surfaces the erosion is, in most cases, from the anode, as it is at inactive surfaces. The arcs are active anode arcs, see Section

2.4(b), which have never been observed at palladium contacts. The metal lost from the anode after a great many active anode arcs tends to be eroded smoothly over the entire surface, like the cathode loss in arcs of the cathode type at palladium contacts. At a moderate pressure of activating vapor, almost all of the metal eroded from a silver anode is transferred to the cathode, but at a high pressure much of it is lost. Whether the metal from the anode is transferred or lost is correlated with the amount of carbon formed by the active arcs; if the production of carbon is small, metal is transferred, but in the presence of much carbon, the metal does not stick to the cathode and is lost. The amount of carbon formed (in air) by active anode type arcs at silver surfaces is very much less than the amount formed by active cathode type arcs at palladium surfaces, and this difference accounts for the fact that a great deal of the eroded metal is transferred at active silver surfaces, although there is always very little transfer at active palladium surfaces.* The erosion of a silver anode by active anode arcs may be as great as 10^{-13} cc/erg, but is lower than this whenever the carbon formation is sufficiently slight to permit much transfer of metal.

Long, sustained break arcs at active silver surfaces become cathode arcs when the electrode separation becomes sufficiently great. Such arcs give cathode erosion resembling that at active palladium surfaces. For a long sustained break arc, the cathode erosion suffered when the electrode separation becomes very large may be greater than the anode erosion occurring when the electrodes are closer together, so that the net loss from the cathode may be the greater. There may even be a small net anode gain.

Measurements of transfer at electrical contacts have sometimes been very confusing in the past, both because of their complexity and because of their apparently erratic character. Now, with well developed insight into the mechanism of short arcs, this complexity of transfer and its varied character have been most useful in improving our understanding of short arcs and of the transfer of metal to which they give rise.

The over-all picture of activation will be given in the following pages.

2. INTERPRETATION OF ACTIVATION

After one has concluded that activation is due to solid carbonaceous material, it is natural that tests should be made upon contacts of solid

* At extremely low pressures of activating vapor, active anode arcs at silver surfaces may not only transfer to the cathode practically all of the metal lost from the anode, but the type of erosion may even be changed to the mound and pit type characteristic of inactive arcs.¹

carbon, and upon metal surfaces on which carbon particles have been dusted. The results of these tests have supplemented measurements upon active noble metal contacts and have led to a great increase in our knowledge of activation. In fact they open the way to a fairly thorough understanding of the subject.

2.1 *Striking Field*

Five different experiments have been carried out, which were designed to discover the reason for the low striking field at active contacts. Although the results of these experiments do not establish the reason for the low striking field in any definitive fashion, they do lead to an explanation which seems entirely satisfying.

The simplest of these experiments has already been reported at the end of Section 1.1. It is the observation that the striking field at active contacts is much the same at different striking voltages, of course below air breakdown only.

In another experiment, not heretofore published, W. S. Boyle and P. Kisliuk produced active spots at various points along a palladium wire. The wire, which lay on the axis of a glass cylinder, was made active at these selected points by repeated short arcs in an atmosphere containing limonene vapor. The other electrode was operated by an electromagnet outside the cylinder, with the magnet arranged so that the electrode could be placed at any location along the wire or withdrawn completely at any time. After activating a number of points, as determined by continuous oscilloscopic observation, the cylinder was exhausted and field emission currents were drawn from the wire to the cylinder. From observation of a fluorescent coating on the inside of the cylinder, it was found that the positions along the wire, which gave the largest currents, were quite unrelated to the active spots. From this experiment, one can conclude that the work function of active spots along the wire was not lower than the work function of other parts of the wire, and also that there was no significant enhancement of field emission at these spots because of roughness. Thus, the activation of contacts by organic vapors is *not* due to enhanced field emission currents because of lowering of the work function or because of greater surface roughness.

In a third experiment by F. E. Haworth,⁷ measurements were made of the electrode separations at which an arc strikes between a palladium electrode and a smooth palladium surface upon which carbon particles had been deposited. For this experiment, solid carbon particles of fairly uniform size were obtained by blowing air at a low controlled rate

TABLE I — EFFECT OF CARBON PARTICLES UPON STRIKING DISTANCE

Range of Particle Size (by microscopic measurement)	Average Striking Distance at 50 Volts	Apparent Striking Field
No Particles	0.10×10^{-4} cm	5×10^6 volts/cm
0 to 1×10^{-4} cm	1.4	0.36
0 to 2.5	2.5	0.20
4 to 5	4.3	0.12

through agitated carbon dust and collecting the particles that had been carried upward for a considerable distance in the air stream. The time of deposition of these particles upon the smooth surface was adjusted to give an average distance between particles of about 10 times their diameters. The smooth palladium surface with a fairly uniform, but sparse covering of carbon particles was made the cathode in measurements of striking distance by the oscilloscope method, Section 1.1. For a particular size of particle, 100 measurements were made of striking distance, each measurement at a different point on the surface, so as not to include any measurement of striking distance at a place on the surface where the original particles had already been burned off.* Table I gives the ranges of particle size as found microscopically and the corresponding average measured values of striking distance. The increase of striking distance was just equal to the particle size. At each arc, a particle was destroyed so that the time to closure measured on the oscilloscope corresponded, not to the true striking distance, but to the distance from the anode to the cathode surface upon which the particle rested. The electric field at which the arc struck was very much higher than the calculated values of the third column of Table I, and was not significantly different from the striking field for inactive surfaces.

In the fourth experiment, the striking field was measured between electrodes of solid carbon. One of these was mounted upon a cantilever bar in such a way that it could be moved through extremely small measured distances by pushing on the end of the cantilever bar using a micrometer screw (Reference 4, page 33). The zero point was found by touching the contacts through a high resistance galvanometer circuit; then the contacts were separated and the striking distance found after applying the voltage. Measurements made in this way by M. M. Atalla (Reference 11, Table I) have given, for the striking field for carbon elec-

* A correct measure of striking distance is obtained only when the arc energy is sufficient to burn up the carbon particles completely. No appreciable mound of metal is thrown up to falsify the distance measurement, because the arcs are of the cathode type, see Section 2.4(a).

trodes, 2.4×10^6 volts/cm, and our unpublished measurements agree with this. The striking field for carbon surfaces is thus only a little less than that found for cathode arcs at clean metal surfaces (Reference 4, Fig. 8), and very different from the field at which arcs strike between active surfaces.*

In another experiment, tests were carried out upon carbon particles in the 4 to 5×10^{-4} cm range of diameters, deposited sparsely upon a palladium surface as before. A careful comparison was made of the electrode separations at which an arc struck at 50 volts and at 250 volts. At the higher voltage, the distance was greater than at the lower voltage by the factor of only 1.3, offering confirmation that the isolated carbon particles act chiefly as chunks of material, partially closing the electrode gap.

The one way in which the carbon that produces activation differs from other carbon, and in particular from small carbon particles dusted sparsely upon a smooth metal surface, is in the very large number of its particles and in its state of subdivision. This gives an eminently plausible clue to the great electrode separation at which breakdown occurs between active surfaces. According to this model, breakdown occurs at a great separation between active surfaces because, at the electric field corresponding to this separation, electrostatic forces become sufficient to cause motion of small particles which decreases the separa-

* In measuring the striking field at carbon surfaces for low voltages by the oscilloscopic method, a value of the order of 0.6×10^6 volts/cm was found earlier (Reference 6, Table I). This result was certainly in error, because of burning of carbon in the arc, so that the separation of the electrodes when the arc ended was greater than it was at the arc initiation.

To check this explanation of the earlier incorrect result, an experiment was carried out in which the time to closure for carbon electrodes was measured as a function of the energy in the arc. In successive tests, a number of different capacitors, each charged to 50 volts, were discharged on the closure of carbon electrodes. The time to closure was found to increase progressively with capacitance for the values 10^2 , 10^3 , 10^4 and 10^5 μmf . Carrying out the measurements many times and taking average values, it was found that the time to closure increased linearly with the cube root of the capacitance. This suggests strongly that a hole was being burned in one of the electrodes and the increased time to closure was just the time for one electrode to move the depth of the hole. A quantitative value for the volume of the hole can be obtained from the data, on the basis of an assumed hole shape. In earlier work (Reference 9, page 1088), a pit on a metal electrode was assumed to be a spherical segment with the depth equal to one-half of the pit radius. Making the same assumption for the hypothetical hole in the present tests, and assuming an electrode velocity on closure of 30 cm/sec, it turns out that the relationship between volume of the hole and energy of the arc is $V = 4.5 \times 10^{-12}$ cm³/erg. The agreement of this result with that for the erosion of the metal anode in an anode arc (Reference 9, page 1088), is remarkable and must be largely fortuitous. The agreement does, nevertheless, make almost certain that burning of one of the electrodes (the cathode, as we know from other work) is the reason for the oscilloscopic method giving incorrect values for the electrode separation at which an arc strikes between carbon electrodes (Reference 6, Table I).

tion. The experimentally observed field of 0.6×10^6 volts/cm is the field at which this motion becomes appreciable for the very small sooty particles. With the start of motion of this sort the field is increased, and further motion is assured making the situation unstable. The gap is greatly decreased in length before *electrical* breakdown takes place, and the field at electrical breakdown is probably as high as it is at any carbon surface.

2.2 Arc Voltage

At the beginning of an active arc, at least one carbon particle is always exploded by the arc current, but only when the surface is very heavily carbonized is an enhanced arc voltage observed, Fig. 4. It must not be thought that the higher arc voltage occasionally found at the beginning of an arc is to be attributed directly to the presence of carbon vapor in the arc during its early stages, because carbon does not have an exceptionally high ionization potential. P. Kisliuk has shown¹² that, in a field emission short arc, the arc voltage should be just slightly larger than the sum of the ionization potential of the metal of the electrodes and of its thermionic work function. Now this result holds quite well for a number of different metal arcs, but does not hold at all for carbon. The short carbon arc is apparently of a different type, and has no well defined arc voltage. On the other hand, although carbon, unlike the noble metals, gives out thermionic electrons copiously long before it is hot enough to vaporize, a true thermionic arc cannot have the enormous current densities that occur in short arcs. (It does not seem impossible that thermionic emission may help to maintain an arc when the current is lower near its end.) The high arc voltage at the beginning and end of an active arc between heavily carbonized surfaces may be due to a dearth of positive ions, requiring a higher applied field to maintain the field emission. In any case, it is like the higher arc voltage of carbon which we do not understand. When the higher arc voltage is not detected, the vaporization of metal must be profuse, and only when vaporization is reduced, as it is when the current is very small near the beginning and end of an arc at the discharge of a capacitor into an inductive circuit, is the higher arc voltage observed. On rare occasions heavily carbonized surfaces show a suddenly enhanced arc voltage for a short interval near the middle of an arc. That this should occur very much less often than at the beginning or end of an arc is understandable.

The observation that the arc voltage sometimes becomes high near the end of an arc suggests strongly that an active arc is moving continually during its life. Only when the current is insufficient to vaporize carbon and underlying metal freely, and thus to maintain the large ion

density necessary for the low voltage field emission arc, does one observe the high and erratic arc voltage characteristic of carbon.

2.3 Minimum Arc Current

Values of minimum arc current for carbon electrodes have already been published. They are of the order of 0.02 to 0.06 ampere and agree fairly well with measurements of minimum arc current for very active metal contacts (Reference 2, Table V). The very low value of the minimum arc current for carbon, either in solid form or dispensed upon the surfaces of active contacts, is related to the low electrical and thermal conductivities of carbon. These low conductivities permit explosion of carbon particles on the cathode by currents too small to vaporize any metal. It has already been pointed out that it is this very low value of minimum arc current which accounts for the greatly enhanced energy that is dissipated at active relay contacts.

From the low value of minimum arc current for active surfaces, one concludes that near its end an active arc is always located at a fresh point on the electrode surfaces, one from which carbon was not burned off earlier in the life of the arc. It had already been concluded from occasional high values of arc voltage near the end of an active arc that this is sometimes true, but the minimum arc current values extend this earlier conclusion to indicate that it is *always* so. An active arc cannot remain in a fixed position as does an inactive anode arc (For example, Reference 4, Fig. 1). The implication is thus suggested that any arc between active palladium contacts is a cathode arc. Further presumptive evidence for this is, of course, furnished by the very much greater electrode separation in the case of active arcs; it is well known¹³ that large distances favor cathode arcs, because at great distances the anode cannot be efficiently heated by electron bombardment.

The interpretation of minimum arc current of active cathode arcs to which we have been led can be written down in words, but we have not succeeded in any quantitative formulation. It is well known that every cathode arc is made up of a great number of small arcs moving continually over the electrode surfaces and exploding one point, or one particle after another on the cathode.^{3, 4} In the case of an *active* arc, the end comes when the current gets so low that it will no longer explode a carbon particle, or when no suitable particles are available.* The much

* This is a necessary criterion for the end of an active arc only in the case of very short arcs. For electrodes that are being pulled apart to break a current larger than the minimum arc current, an arc will, of course, finally fail because of the great electrode separation, even though the current is above the minimum arc current, as in the final failure of the arc in the oscilloscope trace of Fig. 1(b). For inactive anode arcs the minimum arc current arises in a quite different way and has been interpreted in fairly satisfactory quantitative fashion.¹⁴

higher minimum arc currents for cathode arcs at clean surfaces is attributed to the higher thermal and electrical conductivities of metals and to the absence of loose material making poor contact with the surface. This picture is supported by the observation that metal contacts are made temporarily active by almost any kind of loose surface particles of very small size (Reference 2, Page 961).

2.4 Erosion

2.4(a) Palladium and Platinum. Further evidence that an active arc at palladium or platinum surfaces is always a cathode arc is furnished by the fact that the cathode loses much more metal in an active palladium arc than does the anode. (See also Reference 4, Table I).

The direct way of proving that an active arc at palladium or platinum surfaces is a cathode arc would, of course, be microscopic examination of the contact surfaces after a single arc. This is not practicable because surfaces become active only after repeated arcs, but one can do what is apparently quite equivalent by looking at the damage done by a single arc to surfaces on which small carbon particles have been dusted. Experiments by Haworth do indeed prove that arcs at such surfaces are cathode arcs, even at the low striking potential of 50 volts, and when the maximum diameter of the carbon particle is only 1×10^{-4} cm. Fig. 7(b) is typical of many examinations by Haworth of palladium cathodes after a single arc at surfaces upon which carbon particles had been deposited. The striking potential was 50 volts and the capacitance that was discharged was $C = 10^{-8}$ f, so that the energy $C(V_0 - v)v$ was 50 ergs. For comparison, photographs are reproduced in Figs. 7(a) and

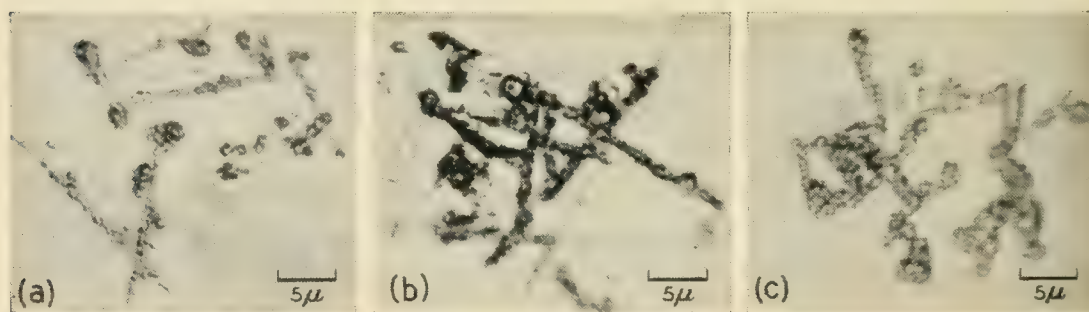


Fig. 7 — Photomicrographs of palladium cathode surfaces after single cathode arcs. The photograph of (b) was obtained after a 50 erg arc with 50 volt striking potential at a surface upon which carbon particles has been deposited. This sort of cathode damage was observed for all of the different sizes of carbon particles which were tested, even for the smallest having diameters of only 10^{-4} cm. The comparison photographs (a) and (c) represent the damage done respectively by 40 erg and 80 erg arcs to palladium surfaces without carbon particles, each arc at the striking potential of 400 volts.

7(c) which show clean palladium cathodes after constant current cathode arcs of 4 amperes lasting, respectively, for 0.072 microsecond and for 0.14 microsecond. The striking potential in each of these arcs was 400 volts, the total arc energy being 40 ergs and 80 ergs. The three photographs of Figs. 7(a), 7(b) and 7(c) represent then the markings made on the cathode by arcs of 40, 50 and 80 ergs respectively. The voltage of 400 was chosen for the two comparison photographs of Figs. 7(a) and 7(c) because this is above the minimum air breakdown potential, and arcs on closure at striking potentials above this value are known to be always cathode arcs (Reference 4, Fig. 4).

Cathode markings such as those of Fig. 7(b) are occasionally produced by arcs at 50 volts on relatively clean palladium surfaces. In general, however, an arc at this low striking potential between clean surfaces is an anode arc, leaving a single well defined pit on the anode, and on the cathode, a single roughened area with considerable metal spattered over from the anode (Reference 9, Fig. 6). While a cathode arc, making on the cathode the type of markings shown in Fig. 7, is rather rare between clean palladium surfaces at a striking voltage as low as 50 (Reference 4, Fig. 4) it is the usual kind of arc between surfaces upon which carbon particles have been dusted, and by implication, it is the sort of arc that occurs between active surfaces. That this arc should cause loss of metal from the cathode is clear from the photographs of Fig. 7, and from the fact that the damage done to the anode sometimes cannot be detected and is always rather slight.*¹⁰ Between clean surfaces, this sort of arc occurs more frequently at higher striking voltages, and invariably on closure when the potential is above the minimum breakdown potential for air. It is the greater striking distance that favors the cathode type of arc, and for active arcs also it is just this enhanced electrode separation, resulting from carbon particles, which can be thought of as the reason for the arc being of this type. There is obviously a critical distance above which arcs are of the cathode type, and for palladium electrodes this critical distance is less than 1×10^{-4} cm. Earlier experiments can be used to define this critical distance better. From the data of Fig. 8 of Reference 4 it appears that this distance for palladium is about 0.5×10^{-4} cm.

Markings made on the cathode by a single arc between active palladium surfaces are doubtless not easily distinguishable from those resulting from a single arc that has been constrained to be of the cathode type only by a high striking potential and the resulting great electrode separation. Nevertheless, when many times repeated, the over-all results

* See the footnote relating to Fig. 3, see page 776.

of cathode arcs between active surfaces, and of cathode arcs between inactive surfaces, are markedly different, as has been pointed out earlier.

The fact that erosion by inactive, or clean-surface, arcs takes the form of a mound on one electrode and a crater in the other means simply that successive arcs tend to occur at the same place on the electrode surfaces. This is because each arc must occur where the electric field between approaching electrodes is highest, and the roughening from one arc will be the site of the highest field before the next discharge occurs.

With carbon particles on the surface, the situation is different. In the case of active cathode arcs, the electric field between approaching electrodes is highest at a point where a group of carbon particles, perhaps pulled up by electrostatic forces, closes a large part of the electrode gap, and an arc must necessarily strike at such carbon particles. In the activating process, carbon is always being formed by an arc, but only at its periphery; at the hottest parts of the arc, carbon which was formed earlier, is completely removed. Not only does each arc move during its lifetime, continually searching out new carbon which was formed earlier, but a later arc will not strike at a point from which carbon was just cleaned by an earlier arc. This restless movement from point to point results finally in erosion that spreads over the surface in a way which is likely to be statistically uniform.

2.4(b) Silver and Gold. Although the character of the erosion at silver (and gold) surfaces, and also its magnitude, are drastically altered by activation, Section 1.4(b), the "direction" of the erosion is still in most cases that characteristic of anode arcs. The predominant loss of metal on closure is usually from the anode for active silver electrodes at voltages too low for air breakdown, just as it is for inactive silver electrodes at low striking voltages. This is in marked contrast to the behavior of palladium surfaces when they become active; for active palladium surfaces, loss of metal is always chiefly from the cathode. The behavior of silver leads naturally to the hypothesis that even when the surfaces are active arcs at low striking voltages are anode arcs, as they are when the surfaces are inactive.

This hypothesis has been subjected to test by F. E. Haworth by the same method used in the case of palladium surfaces. Small carbon particles were dusted on a polished silver surface, and the surface was examined microscopically after it had been subjected to a single arc under the circuit conditions used in similar tests at palladium surfaces. When the maximum particle diameter was 5×10^{-4} cm, it was found from the microscopic examination that all arcs were of the cathode type (see, for example, Fig. 7), but when the maximum diameter was 2.5×10^{-4}

cm all arcs were of the anode type with the characteristic pit on the anode and a roughened spatter of metal on the cathode. It is clear from these tests that active arcs at silver surfaces are of the anode type if the layer of carbonaceous material responsible for activation is not heavy enough to permit an arc to strike at a separation greater than 2.5×10^{-4} cm, but that they are of the cathode type when the layer is sufficiently thick to permit arcs at 5×10^{-4} cm. The electrode separation at which an arc takes place determines the character of the arc.* The critical distance for silver surfaces lies between 2.5 and 5×10^{-4} cm. Erosion at active silver surfaces on closure must be predominantly from the anode unless the layer of activating carbonaceous material is so heavy that arcs strike when the electrode separation is greater than 2.5×10^{-4} cm.

After long continued operation at very high pressures of activating vapor it is sometimes, but not always, found that arcs on closure result in erosion that is chiefly from the cathode. The conclusion drawn from these measurements is that the striking distance at active surfaces on closure at low voltages can sometimes, with considerable difficulty, be made greater than 2.5×10^{-4} cm. Unless great pains are taken to keep surfaces very heavily carbonized, the striking distance on closure at active surfaces at low voltages is of the order of 2.5×10^{-4} cm or less. On closure at voltages that give air breakdown, the erosion of silver is predominantly from the cathode whether the surfaces are active or inactive, because the minimum distance for air breakdown (15×10^{-4} cm) is much above the critical distance for silver.

On breaking active silver contacts in an inductive circuit, erosion is chiefly from the anode unless the arc lasts long enough for the electrode separation to exceed the critical distance of 3 or 4×10^{-4} cm. During the time an arc persists at distances greater than this, the loss is predominantly from the cathode. For velocities typical of a U-type relay, the critical distance may be reached in 10 or 20 microseconds, and equal erosion may be attained in a time of the order of 40 microseconds. If the partial pressure of activating vapor is very high and the surfaces unusually heavily carbonized, much of the eroded metal will be lost. Under more usual conditions of lower vapor pressures, most of the eroded metal is transferred to the opposite electrode. Thus there may be a critical arc duration for which the erosion of each silver electrode is nearly

* Similar tests were carried out upon polished gold surfaces upon which sparse layers of carbon particles had been dusted. For particles of maximum diameter 2.5×10^{-4} cm all arcs were found by microscopic examination of the electrodes to be anode arcs, and for particles in the range of diameters from 4 to 5×10^{-4} cm all arcs were found to be of the cathode type. These results are identical with those found for silver.

zero, and for an arc lasting longer than this time, there may be cathode loss and actual net gain by the anode. No such balancing effect is possible for palladium.

2.4(c) *Anode Arcs and Cathode Arcs.* The model of an active cathode arc to which we have been led seems fairly clear and rather well established, but the model of an active anode arc is more poorly defined. From electron micrographs of the damage done to the cathode by an arc of the cathode type (Reference 4, Fig. 3), it is known that an arc of this type is intermittent, striking over and over again. In an *active* arc of the cathode type, a carbon particle on the cathode is blown up each time the arc strikes, but always there is metal vaporized from the cathode at the site of the particle and the amount of vaporized cathode metal is greater than the amount of vaporized carbon, so that the arc is an arc in metal vapor.

We know less of an active anode arc, and it may well be that some experiments described above seem to imply a model which is not consistent with other observations. The facts that we know are, that at a lightly carbonized silver surface an arc strikes at an electrode separation much greater than the separation at which it would strike if there were no surface carbon, that the resulting arc produces loss of metal predominantly from the anode, and finally that the minimum arc current is very low. The arc is a true anode arc by our implied definition of such an arc, yet it is certainly an active arc. When the arc current is high, a crater is being produced on the anode as in the case of an inactive anode arc, and also in the case of an active anode arc at a surface on which a few carbon particles of diameters not greater than 2.5×10^{-4} cm have been dusted. When the current becomes too low, or is too long sustained, one presumes that the arc is extinguished as in the case of inactive anode arcs.¹⁴ It may then restrike at another carbon particle. One speculates that an anode arc is intermittent when the arc current is very low, being initiated over and over again as are cathode arcs throughout their lives. A carbon particle is exploded repeatedly on the cathode. Yet, because the separation is less than the critical distance, at each re-ignition of the arc, metal vapor is derived from the anode rather than from the cathode, and possibly the over-all anode erosion results in a single anode pit produced when the current was sufficiently high, plus an array of very small anode pits formed while the current was small and intermittent. This model must be regarded as a plausible speculation without support in direct observation. The existence of the active anode arc is well established although the course of such an arc is speculative.

Some insight into the reason for the existence of a critical electrode

separation, determining whether an arc is of the anode type or of the cathode type, can be obtained from a simplified picture of the evaporation of metal in an arc. One assumes a field emission arc just being established between a cathode point and the anode surface, the initial ions being supplied by oxygen and nitrogen of the air with as yet no metal vaporization. The electrons from the point are assumed to travel in straight lines to the anode and cover uniformly an area $\pi(L \tan \theta)^2$ where L is the electrode separation. If i is the total electron current and v the arc voltage, the power density on the anode is $iv/\pi(L \tan \theta)^2$, decreasing with increasing separation. A lower limit for the power put into the cathode point is $(\rho/d)i^2$ where d is the diameter of the point and ρ the resistivity of the cathode metal. Whether the anode begins to vaporize before the cathode, or vice versa, is determined in some way by the ratio of these quantities $BL^2\rho/d$, where parameters unimportant for the present discussion are grouped together in B . For $L^2\rho/d$ greater than some critical value, we shall have cathode evaporation and an ensuing cathode arc, but for $L^2\rho/d$ less than this value, the anode will begin to evaporate first with a resulting anode arc.

The resistivity that probably counts is the resistivity at the melting point. At the temperature of melting, the resistivity of palladium is nine times greater than that of silver. Thus one can expect from this simple model that the critical distance which determines whether an arc is of the cathode or anode type will be three times greater for silver than for palladium. If a silver point is less sharp than a palladium point, d greater for silver than for palladium, as it may be because of the well known property of silver atoms to migrate at room temperature, the factor will be greater than this value of three. Now we have the experimental estimate of 0.5×10^{-4} cm for the critical distance for palladium. This simple theory predicts that the critical distance for silver shall be greater than this by a factor of three, or perhaps more. The experimental critical distance for silver is between 2.5 and 5×10^{-4} cm.

Quantitative measures of the erosion of contacts of palladium and of silver, which were given in Section 1.4, are collected in Table II for ready reference.

From additional experiments, not reported in Section 1.4, it is known that these values of transfer apply approximately for potentials both above and below the minimum breakdown potential for air. Not all types of arcs occur, however, for both palladium and silver at potentials above and below the minimum breakdown potential. At potentials that give air breakdown, all arcs on closure are of the cathode type for both metals whether active or inactive. At potentials that do not give air

TABLE II—LOSS OR GAIN OF METAL FROM ARCING FOR
PALLADIUM OR SILVER
(in units of 10^{-14} cc per erg)

	<i>Cathode</i>	<i>Anode</i>	
Inactive Arcs			
Anode Type.....	4 gain	4 loss	mound and pit
Cathode Type.....	1 loss	1 gain	mound and pit
Active Arcs			
Anode Type.....	(loss)	10 loss*	smooth erosion
Cathode Type.....	4 loss†	(loss)	smooth erosion

* This high figure refers to arcs on closure at very heavily carbonized surfaces; for lightly carbonized silver surfaces the anode loss is less and most of the metal is transferred to the cathode.

† This figure refers to arcs at closure of palladium surfaces. The rate of cathode loss at break of palladium surfaces is significantly less, as pointed out in Section 1.4(a); and in cathode arcs at active silver surfaces the rate of loss is still less.

TABLE III— OCCURRENCE OF DIFFERENT TYPES OF ARCS

	Below Air Breakdown	Air Breakdown
Inactive Arcs		
Anode Type.....	Palladium, Silver	No*
Cathode Type.....	Palladium Only	Palladium, Silver
Active Arcs		
Anode Type.....	Silver Only	No*
Cathode Type.....	Palladium, Silver	Palladium, Silver

* This applies to arcs at closure. Between separating electrodes, air breakdown often occurs when the electrodes are too close together for air breakdown over the shortest path. Under such conditions, arcs between silver surfaces, which are *initiated* by air breakdown, can become anode arcs, and the transfer resulting from such arcs gives dominant anode erosion.

breakdown, all inactive arcs at silver surfaces are of the anode type, and active arcs of the anode type occur for silver only. These facts are tabulated for reference in Table III. All of them are at once predictable from the values of the critical distances for palladium and silver, and from knowledge of the way in which breakdown distance is changed by activation.*

* The difference between the transfer behavior of palladium and silver electrodes in the active condition suggested to R. H. Gumley that the damaging effects of activation can be greatly reduced by constructing a relay with negative contacts of silver and positive contacts of palladium. He tried out this idea and found it to be effective. In the absence of activating vapor, a relay in which the negative contacts are silver and the positive contacts palladium has no merit over a relay in which all contacts are palladium, but when vapor is present, the erosion can, under some circumstances, be much reduced by replacing the negative palladium contacts by silver contacts.

The sort of erosion produced by the different types of arcs is shown in the somewhat conventionalized sketches of Fig. 8. These sketches represent cross-sections of the square mating areas (about 1.3 mm on a side) of heavy type U relay palladium contacts. The contact contours are drawn to scale after the metal transfer resulting from repeated arcs with a total energy of 10^8 ergs, using the values of Table II to convert this energy into volumes of metal. The mounds and pits produced by inactive anode arcs and by inactive cathode arcs are assumed to be spherical segments, each having a height equal to half its radius. The smooth erosion resulting from active arcs would have depths which do not show up at all on the scale of this figure. For each electrode in each of the four cases, the erosion is less than 2 per cent of the total volume of the metal of the contact, and represents a fairly early stage in the expected contact life. The electrode separations at which arcs occur correspond respectively to fields of 8×10^6 , 4×10^6 and 0.5×10^6 volts/cm. The striking voltage is assumed to be 50 and the separations are drawn

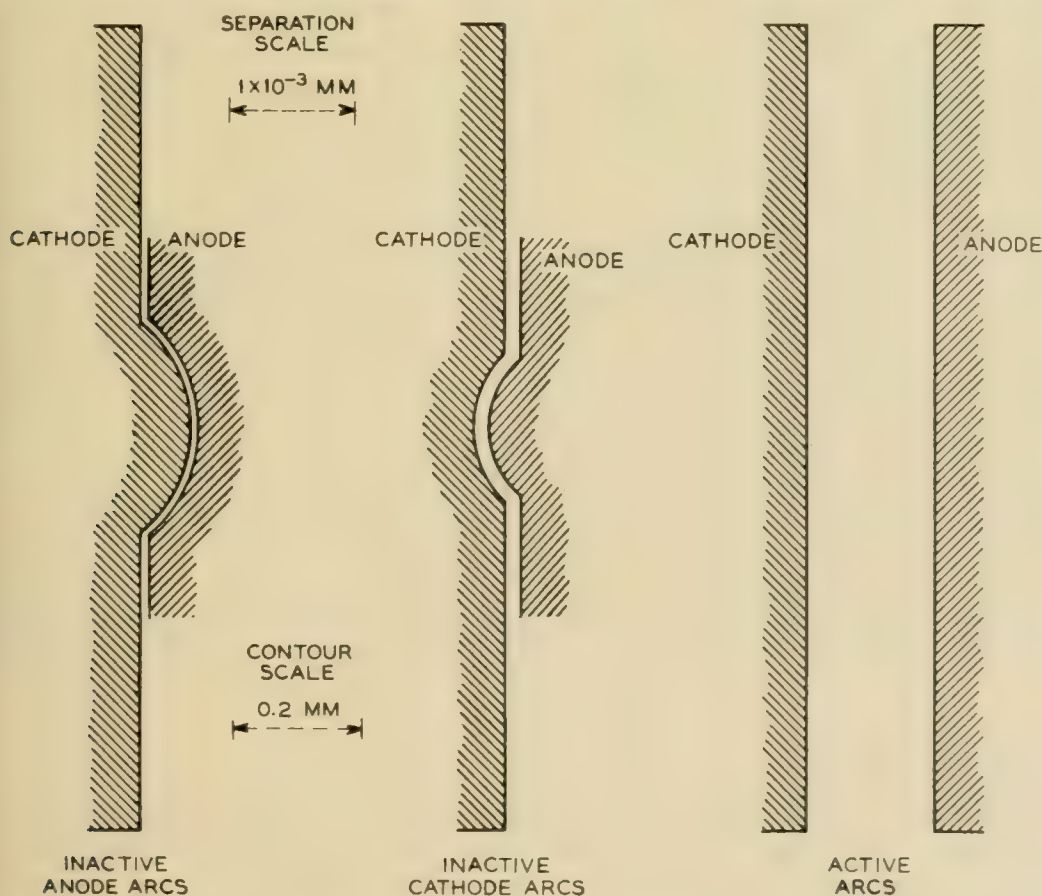


Fig. 8 — Erosion produced by anode arcs at clean surfaces, by inactive cathode arcs and by active arcs of either type, the total energy in each case being 10^8 ergs. The electrode separations at which these arcs strike correspond to 50 volts and are represented here on a greatly expanded scale.

to a scale 200 times greater than the scale of the electrodes. For potentials that give air breakdown, the scale of separation would be changed by large factors.

The sketches of Fig. 8 are of assistance in understanding some of the qualitative erosion differences observed in the four types of arcs (Table II). All of the metal lost from one of the electrodes in an inactive arc of either type comes from the surface of a pit, and from the figure it seems clear that all of it must obviously be intercepted by the other electrode because there is no way for it to escape. This is true even for the case of air breakdown where the electrode separation is much greater ($\sim 15 \times 10^{-4}$ cm). But for active arcs some of the metal coming from each electrode is permanently lost and not transferred to the other side, even though the separation is much less than it is for the case of air breakdown. The permanent loss of metal in the case of active arcs is due to the presence of carbon. When there is carbon on the surfaces, the metal simply does not stick. Chemical analyses have been made of the black powder produced by active cathode arcs at palladium surfaces, and these analyses show palladium metal as well as carbon. The palladium metal lost from the electrodes turns up in this black powder rather than at new locations on the electrodes.

At palladium surfaces, the net loss amounts to most of the eroded metal, but at silver surfaces, most of the metal is transferred. This difference is related to the amount of carbon left on the surfaces. Carbon is found much more abundantly on palladium than on silver, which accounts for the failure of eroded metal to stick to palladium. The greater net carbon production on palladium is due to the low efficiency of cathode arcs (at active palladium surfaces) in burning carbon; the anode arcs, which occur in general at active silver surfaces, are more effective in burning off carbon. It is to be presumed that the amount of organic vapor decomposed per unit of energy at a silver surface is not so very much less than that decomposed at a palladium surface, even though the net carbon left on the surface is tremendously less in the case of silver.

3. RECAPITULATION

We are ready now to state briefly some of the conclusions about active arcs which have been developed above. All of the observations refer to contacts of palladium or of silver. Less extensive tests upon platinum and upon gold have indicated that platinum behaves the same as palladium, and gold the same as silver.

An active arc is an arc that strikes between one electrode and car-

bonaceous material lying upon the other. If one calculates striking field by dividing the potential by the separation between the metal electrodes, a very low value is obtained, but this is the field at which electrostatic forces cause movement of carbon particles to decrease the separation; the true field at which the arc finally strikes between carbonaceous material and the opposing electrode is not significantly lower than the striking field for arcs at clean surfaces. Some or all of the local carbonaceous material is burned up by the arc, and metal vaporized from one of the electrodes is soon fed into the arc so that for most of its life the ions of the arc are metal ions supplied by atoms from one or the other of the electrodes. This is true for even the most heavily carbonized electrodes.

There is a critical electrode separation, characteristic of the metal of the electrodes, which determines whether the arc is an anode type of arc with metal supplied by the anode or an arc of the cathode type with metal supplied by the cathode. If the separation is greater than this critical value the arc is a cathode arc, and less than this value an anode arc. This critical distance is about 0.5×10^{-4} cm for palladium electrodes and of the order of 3 or 4×10^{-4} cm for electrodes of silver. The ratio of these distances is somewhat greater than the ratio of the square roots of the electrical conductivities of the metals at their melting points. The critical distance for palladium is so small that all arcs at active palladium surfaces are cathode arcs. For silver, on the other hand, the critical distance is so large that most arcs at low voltages at silver surfaces are anode arcs. In any practical application of silver electrodes, the carbonaceous material formed is rarely or never in a sufficiently thick layer to result in cathode arcs for closure at low voltages. In the case of separating silver electrodes, an active arc may last until the electrode separation is beyond the critical distance for silver; the erosion occurring after this distance is reached is predominantly from the cathode, and the larger net loss may, on occasion, be from the cathode.

The erosion resulting from repeated arcing at active surfaces is different in character from that produced by inactive arcs. Inactive arcs give rise to a crater on one electrode and a matching mound on the other, with most of the metal from the crater transferred to the mound. Active arcs, on the other hand, produce smooth erosion without craters and mounds, often with considerable net loss of metal which appears mixed with carbon as a black powder. This smooth erosion is accounted for by the striking of each new arc on carbon formed by preceding arcs, together with the burning off of carbon at the center of each arc and the formation of new carbon around its periphery.

PART II — ACTIVATING CARBON

4. COMPOSITION OF ACTIVATING POWDER AND RATE OF PRODUCTION

Experiments have been carried out designed to discover the chemical composition of the carbonaceous material responsible for activation, how much is made per unit of energy in an arc, and where it is made. Now it has been pointed out above that the reason for the uniform erosion in an active arc is the burning off of this black powder by arcs and the consequent continual wandering of successive arcs always to neighboring spots from which the powder has not been burned. This burning off of black powder makes quantitative measurements in air of its rate of formation quite impractical. Activation in vacuum avoids the destruction by burning, and makes possible direct measures of rate of formation; in these tests, the chemical composition of the powder can be found also. All of the quantitative studies of activation in vacuum were made by P. Kisliuk, but the results have not been previously published.

In Kisliuk's experiments two electrodes, which were of platinum, were mounted in a glass chamber so they could be operated by the magnetic field of a coil placed outside the chamber, in the manner of a dry reed switch. The electric circuit was arranged to discharge on each closure a capacitor charged to a fixed voltage, with no current flowing in the circuit as the platinum contacts are separated. Air was pumped out and the contacts operated in benzene vapor at a constant rate, discharging the capacitor a convenient number of times per minute. Every experiment consisted of measuring the pressure in the system, from which was deduced the rate of disappearance of benzene and the rate of evolution of hydrogen resulting from its decomposition, hydrogen being distinguished from benzene by freezing out the latter in liquid nitrogen. The pressures were measured by an RCA thermocouple gauge (1946) which was shown in control tests not to produce benzene decomposition. The benzene, which had been distilled repeatedly to remove water vapor, was used at initial pressures not to exceed 10^{-2} mm Hg determined by a dry ice-acetone bath. The experimental arrangement is shown in Fig. 9.

4.1 *Composition*

In the first experiments with this system it was found, as had been expected, that with continued operation of the contacts in benzene vapor, the pressure rose steadily, although benzene continued to disappear. The pressure changes corresponded to the evolution of 3.2 ± 0.6 molecules of H_2 for each vanishing molecule of benzene, agreeing well with

the theoretical value of 3 for complete decomposition of benzene into carbon and hydrogen. The conclusion from this experiment is that the organic material in the black activating powder is just carbon. The precision allows one to say that, if there is any hydrogen at all left in the black powder, it does not exceed 2 hydrogen atoms for every 15 carbon atoms.

4.2 Rate of Production

In experiments in which the energy in individual arcs was varied, by using different capacitors in the range from 610 μmf to 40,000 μmf and by using the two potentials 58 volts and 232 volts, it was found that a particular arc energy gives the same carbon formation per erg whether the striking voltage is 232 or 58, from which one deduces that formation of carbon depends upon energy rather than upon capacitance or voltage separately. The amount of carbon formed per individual arc increases with the energy of the arc but not so fast as linearly. The experimental values M of amount of carbon formed can be related to arc energy E by the empirical formula $M = KE^{2/3}$, over the range studied from 5 ergs to 1,250 ergs. A tentative explanation of this $\frac{2}{3}$ power relation is given in the next section.

Starting with clean electrodes and measuring the total amount of carbon formed as a function of number of arcs, it was found that the rate of production of carbon is initially low but increases with time, soon

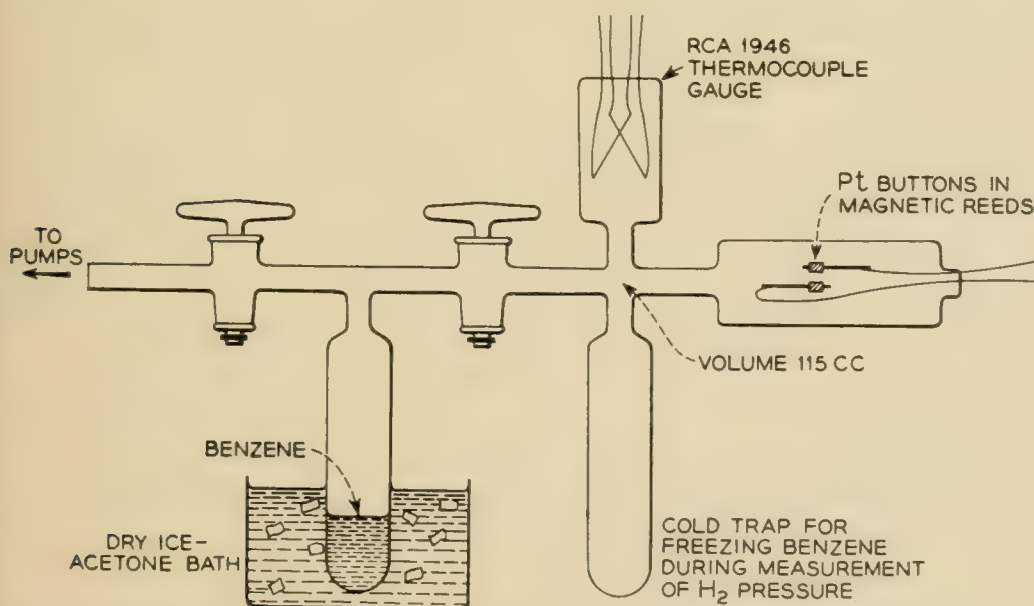


Fig. 9 — Diagram of apparatus used by P. Kisliuk in quantitative measurements of the decomposition of benzene vapor at arcing contacts.

becoming constant. One such set of measurements is plotted as Fig. 10. In this experiment the striking voltage was 232 and the energy in each arc 1,250 ergs. The final slope of the curve of Fig. 10 corresponds to the production of 4.5×10^{-11} gm of carbon per arc — 3.5×10^{-14} gm/erg, 1.8×10^9 atoms/erg — which is 0.04 carbon atom for every electron flowing in the arc, or the decomposition of 3.7×10^{11} benzene molecules per arc, 3×10^8 molecules per erg, or 70×10^{-4} molecule per electron. The lower slope, before the break in the curve, represents the decomposition of 1.1×10^{11} molecules of benzene per arc, or the production of 5×10^8 carbon atoms per erg.

From continuous oscilloscopic observations it was found that the contacts were inactive up to the point where the slope of the curve increased. Here they were slightly active, and beyond this point they were fully active, exhibiting the usual apparent low striking field and low minimum arc current. The amount of carbon required to make the contacts fully active was about 5×10^{-8} gm (2.5×10^{15} atoms) which, if it were in a single spherical speck, would have a diameter of 3.5×10^{-3} cm. Such a speck can be seen quite easily with the naked eye, although an actual deposit of this volume probably could not be seen without a microscope because of its dispersed state.

5. SURFACE ADSORPTION

5.1 *Benzene Molecules on Contact Surfaces*

In an early experiment, the rate of formation of carbon (from measured rate of evolution of H_2) had been found to be independent of benzene pressure down to the lowest pressure tested, which was of the order of 10^{-3} mm Hg. For this reason it was unnecessary to mention absolute pressures in describing the above tests. This lack of dependence on pressure suggests strongly that benzene had been adsorbed on the electrode surfaces and decomposed there, rather than in the space between the electrodes, and that the lowest pressure tested was sufficiently high to keep the surfaces completely covered. This tentative conclusion is confirmed by other considerations given below.

At the pressure of 10^{-2} mm Hg and an electrode separation of 10^{-4} cm, one calculates that only one electron in 3×10^4 can collide with a benzene molecule in the space between the electrodes in the experiment of Fig. 10. The discrepancy between the measured decomposition (70×10^{-4} benzene molecule per electron) and the possible frequency of collision (0.3×10^{-4}) is proof that most of the carbon responsible for activation comes from benzene adsorbed on electrode surfaces rather than from molecules in the space between them.

One gets some insight into the adsorbed films responsible for activation from estimates of the cross-section of an arc and of the amount of benzene adsorbed in a monolayer over an area of this size. A reasonable estimate of the number of molecules in a monolayer of benzene is $7 \times 10^{14} \text{ cm}^{-2}$ (Ref. 16), or $14 \times 10^{14} \text{ cm}^{-2}$ taking into account the two electrodes. Estimates of cross-sectional size have been published for anode arcs, but for cathode arcs the areas are quite different. Since all of the arcs after a surface has become active are certainly cathode arcs, our first concern is with the cross-sectional areas of cathode arcs. It has been observed that the over-all area of the cathode markings made by inactive cathode arcs increases somewhat less rapidly than linearly with total arc energy, and seems to be independent of arc current and arc duration except as they influence the total energy. In one series of experiments, the areas observed (Ref. 10) for low energy arcs corresponded to somewhat less than 10^7 ergs/cm^2 , and to somewhat more than this value for high energy arcs. Assuming for an average value 10^7 ergs/cm^2 , we obtain $1.2 \times 10^{-4} \text{ cm}^2$ for the area of the arcs of the curve of Fig. 10.* This area should have adsorbed on it 1.7×10^{11} benzene molecules. The observed rate of decomposition is 3.0×10^8 benzene molecules per erg or 3.7×10^{11} molecules per arc. Looking at photographs such as those of Fig. 7, one does not feel at all confident that all of the surface in the over-all area of the arc ever became hot enough to decompose benzene. If all of it did become hot enough, the surface must, on the average, have been covered by 2 layers of molecules, and if all of the surface did not become sufficiently hot, by more than two layers. For lower energy arcs, when the number of benzene molecules decomposed per erg is appreciably greater, it is natural to assume that the surface must, on the average, be covered by a still deeper layer of benzene. At least part of the difference between the estimated thicknesses of the layers of benzene molecules for high energy arcs and for low energy arcs can, however, be attributed to the fact that the energy per square centimeter increases with increasing energy, 10^7 ergs/cm^2 being only an average value. The data indicate only that the adsorbed benzene layer is several (greater than 2) molecules thick.

The observed expression $M = KE^{2/3}$ of the above section, relating amount of carbon formed M to total arc energy E , can be accounted for if, in the particular experiment in which this relation was found, the over-all arc area increased with the $\frac{2}{3}$ power of the energy. In various tests

* One should note that the energy density measurements were made upon clean surface or *inactive* cathode arcs, but are being applied here to *active* cathode arcs. Some justification for this is afforded by the fact that the active 50 erg cathode arc of Fig. 7(b) had an over-all area of about $3 \times 10^{-6} \text{ cm}^2$ giving for the energy density $1.5 \times 10^7 \text{ ergs/cm}^2$.

it has been noted that area increases less rapidly than linearly with energy, but it is certain that no universal rule applies in all cases; for example, by restricting the total electrode area, the over-all arc area can be forced to be constant independent of energy.¹⁰ One can conclude only that all of the facts are accounted for by a benzene layer several molecules thick on the electrode surfaces with decomposition by each arc of all of the benzene within its over-all area.

We are now in a position to consider the much lower rate of decomposition of benzene during the initial period before the electrodes became active. In Fig. 10, this lower initial rate is 1.1×10^{11} molecules per arc. This is somewhat less than the number of molecules calculated to lie in a monolayer on the surface covered by an arc. The area used in this calculation was that of a cathode arc, but it is well known that before contacts become active a large proportion of the arcs are anode arcs which have smaller areas (Reference 9, page 1088). The estimated area may, however, be about correct because in the case of an anode arc, carbon is decomposed by heat over an area larger than that of the arc itself.

Within the precision of the estimates we are able to make, it can be said that for inactive contacts operating in benzene vapor each arc decomposes a single layer of adsorbed molecules of benzene. After the contacts become active, the amount decomposed by each arc is greater and is the equivalent of several layers of molecules. It was surmised long ago that much of the vapor adsorbed on active contacts is held by carbon already on the surface rather than by the surface metal. The increased

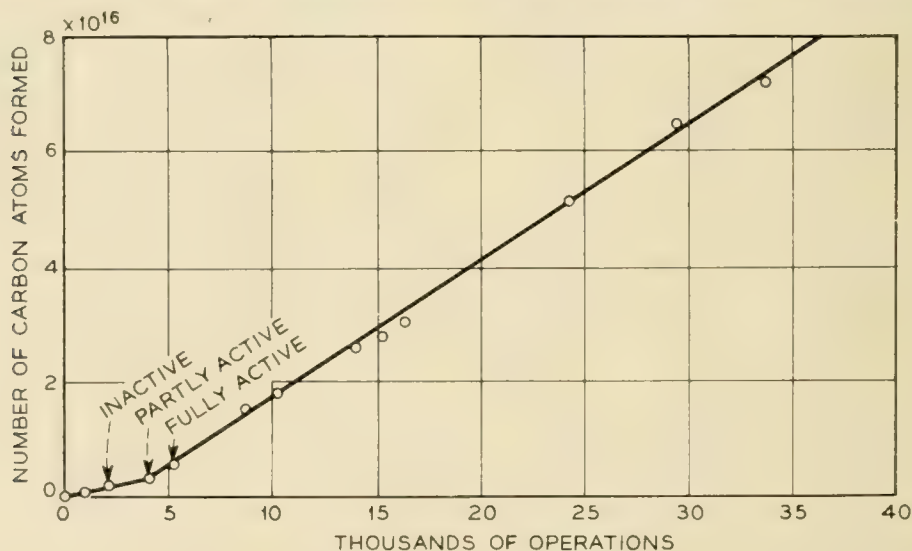


Fig. 10 — Measurements by P. Kisliuk of the amount of carbon formed at arcing platinum contacts, each arc 1,250 ergs. The final slope represents the production of 3.5×10^{-14} gm of carbon per erg of arc energy, or one benzene molecule decomposed for every 150 electrons flowing in the arc.

TABLE IV — BENZENE DECOMPOSITION IN ACTIVE AND INACTIVE ARCS

Measurements		1250 erg arcs (232 volts)	15 erg arcs (58 volts)
Carbon required for full activity	1	2.5×10^{15} atoms	No data
Carbon formed in active arcs	2	1.8×10^9 atoms/erg	7×10^9 atoms/erg
Carbon formed in inactive arcs	3	5×10^8 atoms/erg	No data
Benzene decomposed in active arcs	4	3.0×10^8 molecules/erg	13×10^8 molecules/erg
	5	70×10^{-4} molecule per electron	300×10^{-4} molecule per electron
Calculations			
Absorbed benzene in a monolayer (Ref. 16)	6	14×10^{14} molecules per cm^2	14×10^{14} molecules per cm^2
Benzene molecules struck in space (Compare lines 5 & 7)	7	0.3×10^{-4} molecule per electron at 10^{-2} mm Hg	0.03×10^{-4} molecule per electron at 10^{-3} mm Hg
<i>Active Arcs</i>			
Arc area at 10^7 ergs/ cm^2	8	1.2×10^{-4} cm^2	0.015×10^{-4} cm^2
Number of molecules in one monolayer on arc area	9	1.7×10^{11} molecules	0.02×10^{11} molecules
Decomposed per arc (from line 4)	10	3.7×10^{11} molecules of benzene	0.2×10^{11} molecules of benzene
Effective* thickness of adsorbed layer on basis of 10^7 ergs/ cm^2	11	2.2 molecules	10 molecules
<i>Inactive Arcs</i>			
Arc area	12	$<1.2 \times 10^{-4}$ cm^2	
Number of molecules in one monolayer on arc area	13	$<1.7 \times 10^{11}$ molecules	
Decomposed per arc (from line 3)	14	1.1×10^{11} molecules	No data
Thickness of adsorbed layer	15	1 molecule	No data

* The benzene is probably adsorbed on spongy carbon of much greater true area.

adsorption for contacts already active is doubtless due to the greater surface area resulting from the presence of this carbon.

Many of the numerical values considered here are collected in Table IV for ready reference. These data refer to arcs at platinum surfaces. It is our present opinion that the amount of carbon formed at silver surfaces in similar experiments would be found to be only slightly smaller per unit of energy, although unfortunately no experiments were carried out upon silver.

5.2 Inhibiting Surface Films

One concludes from the above experiments that activation by benzene vapor is the result of firm adsorption of benzene molecules on the elec-

trode surfaces, with heat producing decomposition into carbon and hydrogen rather than evaporation of undamaged molecules. Surface films prevent such strong adsorption, and metals with surfaces that are normally covered by oxide films cannot be activated.

In some very recent experiments in extremely high vacuum, P. Kisliuk has found¹⁷ that benzene molecules are strongly adsorbed upon a tungsten surface that is perfectly clean, but if there is on the surface just one single layer of oxygen molecules, benzene molecules are not adsorbed. M. M. Atalla has reported (Reference 5, page 1090), on the other hand, that tungsten (and nickel also) can be activated if the pressure of air is as low as 10^{-3} mm Hg. It seems probable that arcs at operating contacts remove adsorbed oxygen temporarily, and at sufficiently low air pressures this may be replaced in part by organic molecules rather than by oxygen.

Even at palladium surfaces, some cleaning by arcs seems to be necessary before benzene molecules can be adsorbed. This conclusion is reached in unpublished adsorption experiments carried out by W. S. Boyle upon palladium surfaces in air containing benzene vapor. In this work, two optically flat palladium surfaces are separated by an exceedingly small distance to make an electrical capacitor. With a very sensitive capacitance bridge, one can detect the change in capacity that would be produced by the adsorption on the palladium surfaces of even a small fraction of a monolayer of benzene molecules. In experiments carried out with this equipment it was found that benzene molecules are not adsorbed upon a palladium surface in air at atmospheric pressure. To reconcile this conclusion with the well known facts of activation, it seems necessary to conclude that even a palladium surface can adsorb benzene molecules only after it has been partly cleaned by arcing.

5.3 Alloys

When a base metal is mixed with a noble metal, the result can be an alloy which is activated less readily by organic vapors than would be the noble metal constituent alone. In the curve of Fig. 11 is plotted the number of operations required under a particular set of standard conditions to activate a series of alloys of palladium and nickel. In air, nickel itself cannot be activated at all. The amount of carbon formed from benzene decomposition on the surface of a palladium-nickel alloy is always less than the amount which would be formed under the same conditions upon pure palladium. One does not know whether benzene is held less firmly on the alloy surface so that there is more likelihood

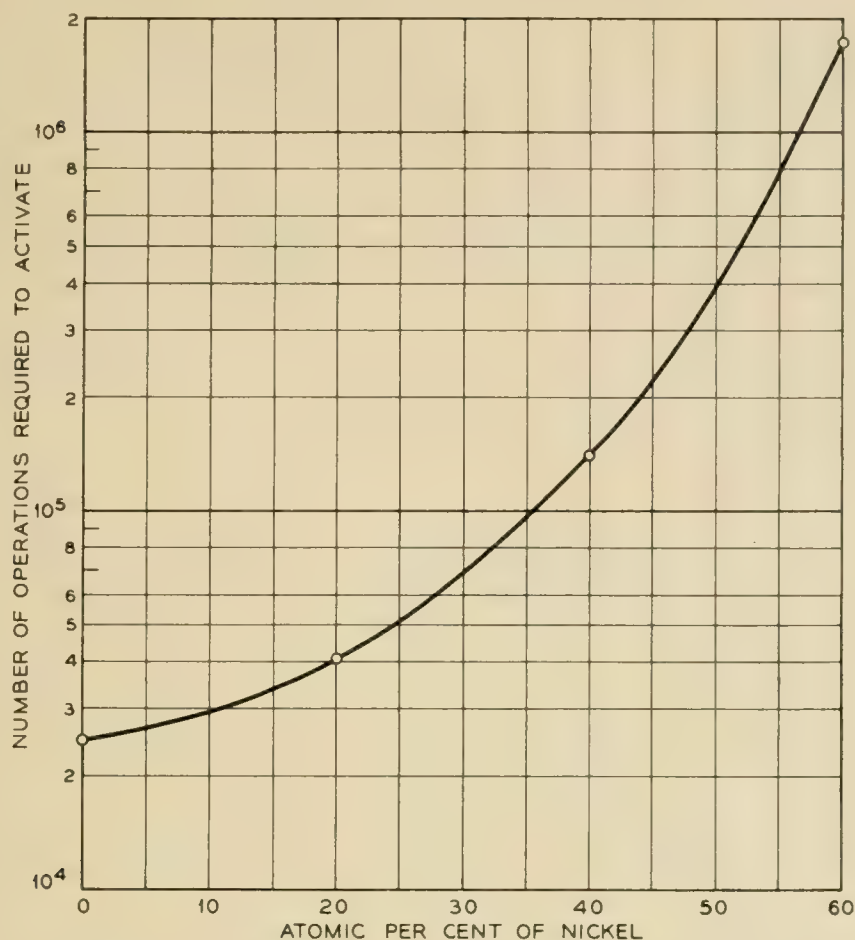


Fig. 11 — Resistance to activation of various alloys of nickel and palladium.

that a heated molecule will evaporate rather than decompose, or whether there are just fewer sites on the surface which can take molecules.

6. ACTIVATION IN AIR

Electrical contacts are not so readily activated by organic vapors in the presence of air as they are when air is absent. Air inhibits the activating process in at least three different ways, and sometimes in a fourth way. These are:

1. Covering up the metal surface so that activating molecules cannot be adsorbed upon it until some of it has been cleaned temporarily by arcing.

2. Offering obstruction in the path of organic molecules on their way to an adsorption site on the metal, so that the molecules diffuse slowly through air up to the surface, whereas in the absence of air, an adsorbed film is formed much more quickly at the same pressure of the organic vapor.

3. Burning off in each arc some of the carbon formed on the surface in preceding arcs.

4. Sputtering and burning off carbon from the cathode in a glow discharge, when such a discharge occurs.

The first of these effects of air has made itself evident in the experiments of Kisliuk, Atalla and Boyle described above. It will not be discussed further. Observations and experiments have been made upon the other three effects of air, and these will be described below. Burning off of carbon in an arc is mentioned first because it can be most nearly separated from other effects and studied individually.

6.1 *Burning of Carbon*

On a surface uniformly covered by organic molecules, carbon must be burned off on the area covered by an arc, but new carbon can be formed on an annular ring surrounding the arc where the metal temperature is lower. As a consequence of this, whereas in the absence of air contacts are activated very much more promptly by high energy arcs than by low energy arcs, in air the situation is less simple. The general result of many experiments is that in air high energy arcs are less efficient in producing activation than are low energy arcs. On the other hand, arcs of extremely low energy are also quite ineffective. There seems to be an optimum arc energy at which contacts can be activated most promptly, which may be of the order of 100 ergs. Activation can be expected to be most prompt when the difference between the area of the arc and the area of the annular ring around the arc is a maximum.

The outer edge of this annular ring is the position on the metal surface at which the maximum temperature just reaches the decomposition temperature of adsorbed organic molecules, about 600°C for the case of benzene. If the width of the ring is Δ and its inner radius R , each arc can be assumed to burn carbon from an area πR^2 , and to form new carbon on an area $\pi[(R + \Delta)^2 - R^2]$. Now Δ certainly increases with increasing energy (being zero for zero energy), but on the simplifying assumption that it is independent of energy, the difference area, which is

$$A = \pi[(R + \Delta)^2 - 2R^2] \quad (1)$$

will be a maximum for that energy that makes R equal to Δ . It is interesting to find the value $R = R_1$ for a 100 erg arc, which is known to be very efficient in producing activation, and then to estimate the maximum temperature reached at the outer edge of the annular ring for $\Delta = R_1$. The simple model predicts that this maximum temperature should be 600°C. When the calculation is carried out in rough fashion, the

temperature is found to be of the order of 300°C , rather than 600°C . The correct order of magnitude gives support to the general ideas behind the theory.

According to this very simple model, activation takes place most promptly for arcs of 100 ergs energy, and for such arcs the *net* carbon formed per arc corresponds to the benzene molecules adsorbed on the area $2\pi R_1^2$, which is obtained from Eq. (1) by setting $R = \Delta = R_1$. In vacuum at the same energy, the carbon formed per arc would come from benzene on the area $4\pi R_1^2$. Thus for 100 erg arcs activation will occur almost as quickly in air as in vacuum, but for arcs of greater energy, much more slowly than in vacuum. Qualitative observation has confirmed this general conclusion.

That this picture is, however, over simplified in a fundamental manner is clear from the effect of electrode contours upon ease of activation. For flat electrodes, activation is very much more prompt when the surfaces make good contact over a large area than when misalignment results in contact on a rather small area. Furthermore, flat contacts can often be activated very promptly under conditions for which crossed wires cannot be activated at all. (Reference 8, page 335). In a qualitative way this is understood, but the inhibiting effect of restricted areas is not amenable to quantitative consideration. This effect makes quite clear that the model of an annular ring about an arc is too idealized to be of much quantitative value.

One might expect that the burning off of carbon would be greatly influenced by atmospheric conditions, and thus the ease of activation would depend upon such conditions. This is indeed found to be the case in experiments in which the air contains water as well as the activating vapor. In unpublished experiments F. E. Haworth determined the number of operations required to activate contacts under a particular set of standard conditions for a wide range of relative humidity. In the range from 10 to 88 per cent relative humidity, the number of operations to make contacts fully active increased exponentially from 1.4×10^3 to 1.0×10^6 , and at relative humidities of 95, 98 and 100 per cent, activation was not attained at all. Furthermore the process of activation could be reversed by water vapor, and contacts that had been made fully active in dry air containing an organic vapor were made completely inactive by continued operation in the same vapor after the addition of water. The effect of water in these experiments may have been due to covering the surfaces so thoroughly with water molecules that the activating vapor could not be adsorbed, or to burning carbon by the water gas reaction, $\text{C} + \text{H}_2\text{O} \rightarrow \text{CO} + \text{H}_2$. The exponential relationship between num-

ber of operations required to make contacts active and relative humidity has no clear interpretation in our present state of knowledge.

6.2 Diffusion of Activating Vapor

In Kisiuk's vacuum experiments the amount of carbon formed and the degree of activation attained was independent of benzene vapor pressure. This is not at all the case when activation is produced by operating contacts in air. In fact, one of the earliest observations was a minimum vapor pressure below which contacts could not be activated (Reference 2, Table I). In more careful later tests it was found that the minimum vapor pressure is a function of rate of operation of the contacts, the minimum pressure being actually proportional to the rate of operation over a factor of 100 which was the range tested (Reference 8, Fig. 2). Obstruction offered by air supplies the explanation of this rate effect. Activation cannot occur if electrodes are separated between one arc and the next for a time which is short in comparison with the time required to cover the surface with one monolayer of organic molecules. A rough order of magnitude calculation confirms this conclusion.

As an approximation, one assumes one dimensional diffusion to an electrode surface from the space in front of it, with all molecules reaching the surface sticking to it. Boundary conditions for the solution of the diffusion equation, $\partial C/\partial t = D\partial^2 C/\partial x^2$, are then:

$$C = C_0 \text{ at } t = 0 \text{ for } x > 0$$

$$C = 0 \text{ at } x = 0 \text{ for all values of } t$$

The concentration of activating molecules in the space in front of the electrode is then $C = C_0 \operatorname{erf} [x/2 (Dt)^{1/2}]$. The total number of molecules to have reached the surface at any time t_1 is,

$$m = \int_0^{t_1} D \left. \frac{\partial C}{\partial x} \right|_{x=0} dt = 2C_0(D/\pi)^{1/2} t_1^{1/2}$$

expressed in molecules/cm², when C_0 is given in molecules/cm³. We are interested in the value of t_1 for which m is the number of molecules in a monolayer, and the maximum rate of operation of contacts for activation to occur can be expected to be comparable with

$$n = \frac{1}{2}t_1 = 2C_0^2 D/\pi m^2 = 8.1 \times 10^{32} D(p/m)^2, \quad (2)$$

where p is the partial pressure of activating vapor in mm Hg. The factor $\frac{1}{2}$ in $n = \frac{1}{2}t_1$ appears because diffusion to the surface can occur only when the electrodes are separated, and it is assumed that they are separated for half of the time.

The best data we have for testing this relation are represented by ex-

periments upon activation in vapor of the organic compound fluorene.⁸ According to the observations, the critical rate of operation was found to be proportional to the partial pressure of fluorene rather than to its square as in (2). This is a discrepancy which must be overlooked in our present state of knowledge. To test (2) for fluorene at 20°C, we require values of D , the diffusion coefficient of fluorene in air, p , the partial pressure of fluorene at 20°C, and m , the number of adsorbed fluorene molecules per cm² of surface. The value of $D = 0.067$ cm²/sec. was obtained from a linear relation between $1/D$ and (molecular weight)^{1/2}, which holds quite well for a number of organic compounds. The value $p = 0.04$ mm Hg is the geometrical mean between 0.23 and 0.007 mm Hg, respectively the vapor pressures of naphthalene and anthracene at 20°C. We have estimated $m = 3.3 \times 10^{14}$ molecules/cm², which is related to the corresponding number for benzene, 7×10^{14} in the inverse ratio of the molecular weights.¹⁶ These numerical values give from (2)

$$n = 0.75 \text{ operation/second}$$

as the critical rate that will just permit one monolayer in the time the contacts are separated. The observed critical rate for activation at 20°C from Fig. 2 of Reference 8 is 3. The agreement is pretty good when the crudeness of the model is considered.

6.3 *Sputtering and Burning in a Glow Discharge*

If both arcs and glow discharges occur when electrical contacts are operated in an atmosphere containing an activating organic vapor, the activation of the contacts resulting from the arcs is inhibited by the occurrence of the glow discharges.* This effect is sometimes very beneficial in extending the life of telephone relay contacts. In fact a very simple protective network, consisting only of an inductance of the order of 10⁻⁴ henry placed very close to one of the contacts, has been devised which, under some conditions, will increase the contact life by a factor of about 10.

Quantitative measurements have been made of this inhibiting action of a glow discharge, and from them it has been concluded that the effect is attributable to sputtering and burning of carbon in the discharge. In making these measurements, a pair of contacts was operated in an atmosphere containing limonene vapor in such a way that arcs and glow discharges occurred alternately in controlled fashion. A charged capacitor was discharged in an arc at each closure. By the use of an auxiliary synchronized relay in series with one of the contacts, the circuit was

* It should be pointed out incidentally that a glow discharge in air can also activate silver electrodes. It produces silver nitrite on their surfaces,¹⁸ and silver electrodes with a layer of nitrite are fully active until the layer is burned off.

changed periodically so that a glow discharge could be made to occur at each contact break, or at every 6th, 60th, or 600th break. The glow current was always 0.04 ampere lasting for a time that could be accurately set by means of a synchronized shunt tube.

In all of the tests, the energy in each closure arc was 190 ergs. Measurements were made at partial pressures of limonene of 0.05 and 1 mm Hg. At the lower pressure it was found that the contacts remained inactive indefinitely whenever the time of glow discharge on break was on the average more than 0.25 microsecond for each closure arc, and activation would ultimately take place if the average glow time per closure arc was less than this value. (At the limonene pressure of 1 mm Hg there was a corresponding critical glow time of about 1 microsecond). The obvious interpretation of these tests is that a glow discharge of 0.04 ampere lasting for 0.25 microsecond sputters and burns off as much carbon as is made by an arc of 190 ergs under the conditions of the experiment. To test this conclusion, one needs to know how much carbon is produced by an arc of 190 ergs, and one needs to know the sputtering rate of carbon in a normal glow in air at atmospheric pressure.

Measurements of the sputtering of carbon in a normal glow discharge were undertaken by F. E. Haworth, since such data are not available in published literature. Carbon and graphite electrodes were weighed before and after a normal glow discharge of 0.006 ampere lasting for various lengths of time. The loss of the carbon or graphite negative electrode in nitrogen was found to amount to about 0.15 atom per ion of the discharge. In air the loss was much greater, four times larger for graphite and 15 times larger for carbon (2.3 carbon atoms/ion for carbon in air). The increase in air was attributed to burning, and the difference between carbon and graphite losses in air was believed to be due to smaller crystal size and looser bonding in the carbon case.*

If we use the highest loss figure of 2.3 carbon atoms/ion we find that a glow discharge of 0.04 ampere for 0.25 microsecond should remove 14×10^{10} carbon atoms. From Table IV, one finds that a 190 erg inactive arc in activating benzene vapor produces 9.5×10^{10} carbon atoms in the absence of air (line 3), and an active arc produces 34×10^{10} carbon atoms (line 2).† The net carbon which is left after each arc in air is, of course, considerably less than it would be in the absence of air (Section 6.1), but the order of magnitude agreement between these numerical

* The sputtering rate of 0.15 atom/ion for carbon in nitrogen in the normal glow is about what is reported by Güntherschulze¹⁹ for silver in the *abnormal* glow but is greater by a factor of about 400 than that found for silver in the *normal* glow in experiments by Haworth.¹⁸ Obviously sputtering rates for carbon are exceptionally high.

† It is believed that these figures are substantially the same for limonene and for benzene.

values leaves little doubt that we have correctly interpreted the inhibiting effect of glow discharges upon activation.

6.4 "*Hysteresis*" Effects

Sometimes a pair of completely inactive electrical contacts of a noble metal can be activated very quickly, and an apparently identical pair of contacts cannot be activated at all under exactly the same experimental conditions. In the first case a great amount of carbon may be formed, and in the second case no detectable carbon at all. In order to clear up this confusion, some controlled experiments were carried out upon the activation and deactivation of silver and palladium electrodes in air containing benzene vapor at various partial pressures. From these experiments, it has been possible to relate the variability of earlier results to previous history of the contacts, and the entire behavior is now quite well understood.

In these tests adjustable benzene vapor pressure was obtained by first bubbling air at a controlled rate through benzene maintained at constant temperature by a bath of acetone and dry ice, and then mixing the saturated air with clean air in the proper proportions. In certain tests silver contacts were operated in air flowing from this apparatus, discharging a capacitor on each closure. The number of operations required to produce complete activation was measured for many different values of the benzene vapor pressure. With contacts that had been cleaned in a standard way before each test, it was found that the number of closures required for activation rose extremely rapidly with decreasing vapor pressure over a narrow range of pressures. There was always a lower pressure limit below which it seemed impossible to activate the contacts at all. All of the tests were made at the high operating rate of 60 closures per second.

When the contacts had been cleaned by abrasion before each individual test, the minimum pressure below which activation could not be attained was of the order of 2 mm Hg. (This pressure was exceptionally high because of the high operating rate, see Section 6.2.) A different result was found for contacts that had been previously activated and then cleaned only by repeated arcing; for these contacts the minimum pressure for activation was about 0.7 mm Hg. The factor of 3 between these minimum pressures is doubtless related to the fact that, for those electrodes which had been cleaned by arcing only, there existed neighboring carbonized areas which were never cleaned. Each such area can be expected to hold about three times as much adsorbed benzene on the average as does the same area of clean metal, see Section 5.1, and especially lines 2 and 3 of Table IV. Thus for such surfaces more carbon can

be expected to be formed by each arc. The model is not sufficiently well defined to permit any more exact conclusions.

In another experiment, silver electrodes, which had first been completely activated, were operated for a long period at a greatly reduced benzene vapor pressure. It was found that they remained completely active unless the pressure was very much less than the minimum of 0.7 mm Hg at which activation could be produced. In repeated tests at a variety of low benzene vapor pressures, the number of operations required for the contacts to become inactive was recorded. This number was found to increase very abruptly with increasing vapor pressure, and above about 0.02 mm Hg the contacts remained active indefinitely. This result must again be related to the capacity of a mass of spongy carbon to hold a great amount of adsorbed benzene. Very probably the upper limit of pressure below which contacts cannot be deactivated, depends upon the thickness of the carbon layer produced before the benzene pressure is lowered.

Similar but less extensive experiments were carried out with palladium electrodes.

These hysteresis effects observed in the activation and deactivation of contacts seem capable of explaining the erratic observations that had been made previously. If the immediate history of contacts is sufficiently well known, behavior can perhaps be predicted fairly well for various experimental conditions.

7. BROWN DEPOSIT

Closely related to the activation of relay contacts is the formation of polymerized layers of organic material upon contact surfaces as a result of friction. This material, which is commonly known as "brown deposit", is produced at contacts which do not make or break current. Its mode of formation is thus entirely different from that of the carbon which is the cause of activation. Both have, however, a common origin in layers of organic molecules adsorbed upon surfaces. Discovery of brown deposit and most of the investigation of it were carried out elsewhere (Ref. 20), but some discussion of brown deposit is appropriate here because of its relation to the carbon of activation and because of a study of its formation by P. Kisliuk.

7.1 *Composition*

The composition of brown deposit was determined by Kisliuk in an apparatus similar to that used to investigate the carbonaceous material responsible for activation, Fig. 9, and by the same analytical procedure. The apparatus was modified so that a palladium or platinum electrode

could be rubbed back and forth upon another electrode of the same material. The driving force was a magnet outside the glass apparatus.

In tests carried out in benzene vapor in the absence of air, it was found that the deposit formed on the electrodes contained 65 per cent as much hydrogen as was in the original benzene, about 2 atoms of hydrogen for every 3 carbon atoms, this figure having a possible experimental error of as much as 20 per cent. The brown deposit formed by friction thus differs significantly from the pure carbon produced by arcing which is responsible for activation.* The experimentally determined composition of the brown deposit does not, of course, distinguish between hydrogen or benzene simply adsorbed in the deposit and hydrogen existing in it in some combined form.

7.2 Rate of Production

In Kisliuk's experiments, which were carried out in the absence of air, the rate of production of brown deposit was found to be independent of benzene vapor pressure down to 3×10^{-3} mm Hg, which was the lowest pressure tested, just as was the case in the formation of carbon by arcs.

When air is present, the rate of formation of brown deposit may depend upon vapor pressure of the organic molecules. Unpublished experiments have indicated, furthermore, that there may be a limiting vapor pressure below which the deposit does not form, with this pressure dependent upon the idle period between operations.²⁰

In some of Kisliuk's vacuum tests a palladium electrode was rubbed back and forth over an area determined microscopically to be about 4×10^{-3} cm², and produced the polymerization on each rub of 2.1×10^{10} molecules of benzene, or 5×10^{12} molecules per cm² of rub. This is smaller than the number of molecules in a monolayer (7×10^{14} per cm², Reference 16) by a factor of 140. Part of the discrepancy is certainly due to the fact that the true area of contact of the electrodes is less than the apparent area as seen under the microscope. From more careful estimates of area it has been found by other observers that the amount of benzene that is polymerized by friction is, in general, comparable with that adsorbed as a monolayer on the rubbing surfaces.

7.3 Brown Deposit and the Carbon of Activation

Although both brown deposit and the carbon of activation are produced from the decomposition of adsorbed organic molecules, there are

* In this connection, it is interesting to point out, however, that any metal surface, upon which brown deposit has been produced by friction in an appropriate atmosphere, is found to be fully active when tested in a suitable circuit. This activity naturally does not last after the brown deposit has been burned off. In this characteristic, the brown deposit behaves like any foreign more or less insulating layer upon a contact surface.

several differences in the conditions necessary for formation. The carbon is produced on a noble metal but not on a base metal (in air); brown deposit, on the other hand, has been formed on vanadium, molybdenum and tantalum, but it has never been produced on silver and only sparingly on gold.²⁰ The failure of electrodes of silver and of gold to form brown deposit has been associated with the high thermal conductivities of these metals, with the idea in mind that polymerization of organic molecules to brown deposit requires frictional heat. Whether this is true has not been established.

Both brown deposit and the carbon of activation can be formed from any of a great variety of unsaturated ring compounds. Various unsaturated aliphatic compounds which have been tested, and some saturated aliphatic compounds (for example, pentane), can be made to produce brown deposit to a limited extent, but activation has never been attained with any aliphatic compound. It seems probable that some activating carbon is produced from these compounds but the burning off in the arc makes activation impossible.

ACKNOWLEDGMENT

The work reported here is the joint effort of a number of persons whose contributions are acknowledged in the appropriate places. The authors coordinated the investigation and are responsible for its general plan. The authors are indebted furthermore to R. H. Gumley for many helpful criticisms.

REFERENCES

1. R. H. Gumley, Bell Lab. Record, **32**, p. 226, 1954.
2. L. H. Germer, J. Appl. Phys., **22**, p. 955, 1951.
3. L. H. Germer and W. S. Boyle, Nature, **176**, p. 1019, 1955.
4. L. H. Germer and W. S. Boyle, J. Appl. Phys. **27**, p. 32, 1956.
5. M. M. Atalla, B.S.T.J., **34**, p. 1081, Sept., 1955.
6. L. H. Germer, J. Appl. Phys. **22**, p. 1133, 1951.
7. F. E. Haworth, J. Appl. Phys., **28**, p. 381, 1957.
8. L. H. Germer, J. Appl. Phys. **25**, p. 332, 1954.
9. L. H. Germer and F. E. Haworth, J. Appl. Phys., **20**, p. 1085, 1949.
10. L. H. Germer, to be published.
11. M. M. Atalla, B.S.T.J., **32**, p. 1493, Nov., 1953.
12. P. Kisliuk, J. Appl. Phys., **25**, p. 897, 1954.
13. L. H. Germer, Electrical Breakdown between Close Electrodes in Air, to be published.
14. W. S. Boyle and L. H. Germer, J. Appl. Phys., **26**, p. 571, 1955.
15. J. J. Lander and L. H. Germer, J. Appl. Phys., **19**, p. 910, 1948.
16. B. M. W. Trapnell, Advances in Catalysis, Vol. III, Academic Press, New York, 1951, pp. 1-24.
17. P. Kisliuk, to be published.
18. F. E. Haworth, J. Appl. Phys., **22**, p. 606, 1951.
19. A. Güntherschulze, Z. Physik, **36**, p. 563, 1926.
20. H. W. Hermance and T. F. Egan, 1956 Electronics Symposium, Electrical Engineering, to be published.

Bell System Technical Papers Not Published in this Journal

AARON, M. R.¹

The Use of Least Squares in Network Design, Trans. I.R.E., PGCT, CT-3, pp. 224-231, Dec., 1956.

ANDERSON, J. R.¹

A New Type of Ferroelectric Shift Register, Trans. I.R.E., PGEC, EC-5, pp. 184-191, Dec., 1956.

ANDERSON, P. W., see Clogston, A. M.

ANDREATCH, P., JR.,¹ and THURSTON, R. N.¹

Disk-Loaded Torsional Wave Delay Line. I — Construction and Test, J. Acous. Soc. Am., 29, pp. 16-19, Jan., 1957.

AUGUSTYNIAK, W. M., see Wertheim, G. K.

BALA, V. B., see Matthias, B. T.

BASHKOW, T. R.,¹ and DESOER, C. A.¹

A Network Proof of a Theorem on Hurwitz Polynomials and its Generalization, Quarterly Appl. Math., 14, pp. 423-426, Jan., 1957.

BOND, W. L., see McSkimin, H. J.

BOZORTH, R. M.¹

Magnetic Properties of Materials, Am. Inst. Phys. Handbook, Chapter 5, pp. 206-244, Feb., 1957.

BREIDT, P., JR.,¹ GREINER, E. S., and ELLIS, W. C.¹

Dislocations in Plastically Indented Germanium, Acta Met., Letter to the Editor, 5, p. 60, Jan., 1957.

¹ Bell Telephone Laboratories.

BRIDGERS, H. E., see tabulation at the end.

BRIDGERS, H. E., see tabulation at the end.

BURKE, P. J.

. **The Output of a Queueing System**, Operations Research, **4**, pp. 699–704, Dec., 1956.

CLEMENCY, W. F.,¹ ROMANOW, F. F.,¹ and ROSE, A. F.²

The Bell System Speakerphone, Elec. Engg., **76**, pp. 189–194, March, 1957.

CLOGSTON, A. M.,¹ SUHL, H.,¹ WALKER, L. R.,¹ and ANDERSON, P. W.¹

Ferromagnetic Resonance Line Width in Insulating Materials, J. Phys. Chem. Solids, **1**, pp. 129–136, Nov., 1956.

CORENZWIT, E., see Matthias, B. T.

D'AMICO, C.,¹ and HAGSTRUM, H. D.¹

An Improvement in the Use of the Porcelain Rod Gas Leak, Rev. Sci. Instr., **28**, p. 60, Jan., 1957.

DESOER, C. A., see Bashkow, T. R.

DILLON, J. F., JR.¹

Ferrimagnetic Resonance in Yttrium Iron Garnet, Phys. Rev., Letter to the Editor, **105**, pp. 759–760, Jan. 15, 1957.

DITZENBERGER, J. A., see Fuller, C. S.

DOBA, S., JR.¹

The Measurement and Specification of Nonlinear Amplitude Response Characteristics in Television, Proc. I.R.E., **45**, pp. 161–165, Feb., 1957.

EDELSON, D., see tabulation at the end.

¹ Bell Telephone Laboratories.

² American Telephone and Telegraph Company.

ELLIS, W. C., see Breidt, P., Jr.

EVANS, D. H.¹

A Positioning Servomechanism With A Finite Time Delay and A Signal Limiter, Trans. I.R.E., PGAC, AC-2, pp. 17-28, Feb., 1957.

FLASCHEN, S. S., see tabulation at the end.

FLASCHEN, S. S., see Garn, P. D.

FRY, T. C.¹

Automatic Computer in Industry, Am. Stat. Assoc. J., **51**, pp. 565-575, Dec., 1956.

FULLER, C. S.,¹ and DITZENBERGER, J. A.¹

Effect of Structural Defects in Germanium on the Diffusion and Acceptor Behavior of Copper, J. Appl. Phys., **28**, pp. 40-48, Jan., 1957.

FULLER, C. S.,¹ and MORIN, F. J.¹

Diffusion and Electrical Behavior of Zinc in Silicon, Phys. Rev., **105**, pp. 379-384, Jan. 15, 1957.

FULLER, C. S., see Reiss, H.

GALT, J. K.,¹ AND KITTEL, C.⁴

Ferromagnetic Domain Theory, Solid State Physics: Advances in Research and Applications (book), 3, pp. 437-564, 1956, Academic Press, Inc., New York.

GARN, P. D.,¹ and FLASCHEN, S. S.¹

Analytical Applications of Differential Thermal Analysis Apparatus, Anal. Chem., **29**, pp. 271-275, Feb., 1957.

GARN, P. D., and FLASCHEN, S. S.

Detection of Polymorphic Phase Transformations by Continuous Measurement of Electrical Resistance, Anal. Chem., **29**, pp. 268-271, Feb., 1957.

¹ Bell Telephone Laboratories.

⁴ University of California, Berkeley.

GARN, P. D., see tabulation at the end.

GAST, R. W.⁵

Field Experience with the A2A Video System, Elec. Engg., **76**, pp. 44-49, Jan., 1957.

GEBALLE, T. H., see Kunzler, J. E.

GIBBONS, J. F.¹

A Simplified Procedure for Finding Fourier Coefficients, Proc. I.R.E., Letter to the Editor, **45**, p. 243, Feb., 1957.

GREINER, E. S., see Breidt, P., Jr.

GROSS, W. A.¹

The Second Fundamental Problem of Elasticity Applied to a Plane Circular Ring, J. Appl. Math. and Phys., **8**, pp. 71-73, 1957.

GOULD, H. L. B., and WENNY, D. H.¹

Supermendur — A New Rectangular-Loop Magnetic Material, Elec. Engg., **76**, pp. 208-211, March, 1957.

HAGSTRUM, H. D.¹

Effect of Monolayer Adsorption on the Ejection of Electrons from Metals by Ions, Phys. Rev., **104**, pp. 1516-1527, Dec. 15, 1956.

HAGSTRUM, H. D., see D'Amico, C.

HAMMING, R. W.¹

Harnessing the Digital Computers, Columbia Engg. Quarterly, **10**, pp. 16-19, 54, March, 1957.

HANSON, R. L.,¹ and KOCK, W. E.¹

Interesting Effect Produced by Two Loudspeakers Under Free Space Conditions, J. Acous. Soc. Am. Letter to the Editor, **29**, p. 145, Jan., 1957.

¹ Bell Telephone Laboratories.

⁵ New York Telephone Company.

HAWKINS, W. L., see tabulation at the end.

HERRING, C.¹

Theoretical Ideas Pertaining to Traps or Centers, Photoconductivity Conference (book), pp. 81-110, 1956. John Wiley & Sons, New York.

HULL, G. W., see Kunzler, J. E.

KARP, A.¹

Japanese Technical Captions, Proc. I.R.E., Letter to the Editor, 45, p. 93, Jan., 1957.

KING, B. G.¹

Discussion on "An Investigation into Some Fundamental Properties of Strip Transmission Lines with the Aid of an Electrolytic Tank", Proc. I.R.E., 104, p. 72, Jan., 1957.

KITTEL, C., see Galt, J. K.

KOCK, W. E., see Hanson, R. L.

KOWALCHIK, M., see Thurmond, C. D.

KUNZLER, J. E.,¹ GEBALLE, T. H.,¹ and HULL, G. W.¹

Germanium Resistance Thermometers Suitable for Low-Temperature Calorimetry, Rev. Sci. Instr., 28, pp. 96-98, Feb., 1957.

KULKE, B.,¹ and MILLER, S. L.¹

Accurate Measurement of Emitter and Collector Series Resistances in Transistors, Proc. I.R.E., Letter to the Editor, 45, p. 90, Jan., 1957.

KUNZLER, J. E., see tabulation at the end.

LANDER, J. J., see Thomas, D. G.

LAW, J. T., see tabulation at the end.

¹ Bell Telephone Laboratories.

LUNDBERG, C. V., see tabulation at the end.

LUNDBERG, J. L., see tabulation at the end.

MATTHIAS, B. T.,¹ WOOD, E. A.,¹ CORENZWIT, E.,¹ and BALA, V. B.¹

Superconductivity and Electron Concentration, J. Phys. Chem. Solids, **1**, pp. 188–190, Nov., 1956.

McMILLAN, B.¹

Two Inequalities Implied by Unique Decipherability, Trans. I.R.E., PGIT, IT-2, pp. 115–116, Dec., 1956.

McSKIMIN, J. H.,¹ and BOND, W. L.¹

Elastic Moduli of Diamond, Phys. Rev., **105**, pp. 116–121, Jan. 1, 1957.

MENDIZZA, A.¹

The Standard Salt Spray Test — Is It A Valid Acceptance Test?, Plating, **44**, pp. 166–175, Feb., 1957.

MENDIZZA, A.¹

The Standard Salt Spray Test — Is It A Valid Acceptance Test?, Symposium on Properties, Tests and Performance of Electrodeposited Metallic Coatings, A.S.T.M. Special Tech. Publication 197, pp. 107–117, 1957.

MILLER, R. C., see Smits, F. M.

MILLER, S. L., see Kulke, B.

MORIN, F. J., see Fuller, C. S.

NELSON, L. S., see Tabulation at the end.

PALMQUIST, T. F.⁶

Multiunit Neutralizing Transformers, Elec. Engg., **76**, p. 201, March, 1957.

¹ Bell Telephone Laboratories.

⁶ Bell Telephone Company of Canada, Montreal.

PATERSON, E. G. D.¹

Some Observations on Quality Assurance and Reliability, Proc. Third National Symp. on Reliability and Quality Control in Electronics, pp. 129-132, Jan., 1957.

PIETRUSZKIEWICZ, A. J., see Reiss, H.

POTTER, J. F., see tabulation at the end.

PRINCE, E.,¹ and TREUTING, R. G.¹

The Structure of Tetragonal Copper Ferrite, Acta Crys., **9**, pp. 1025-1028, Dec., 1956.

READ, M. H., see Van Uitert, L. G.

READ, W. T., JR.¹

Dislocation Theory of Plastic Bending, Acta Met., **5**, pp. 83-88, Feb., 1957.

REISS, H.¹

Theory of the Ionization of Hydrogen and Lithium in Silicon and Germanium, J. Chem. Phys., **25**, pp. 681-686, Oct., 1956.

REISS, H., see tabulation at the end.

REISS, H.,¹ FULLER, C. S.,¹ and PIETRUSZKIEWICZ, A. J.¹

Solubility of Lithium in Doped and Undoped Silicon, Evidence for Compound Formation, J. Chem. Phys., **25**, pp. 650-655, Oct., 1956.

ROMANOW, F. F., see Clemency, W. F.

ROSE, A. F., see Clemency, W. F.

ROSE, D. J.¹

Microplasmas in Silicon, Phys. Rev., **105**, pp. 413-418, Jan. 15, 1957.

SCHLABACH, T. D., see tabulations at the end.

¹ Bell Telephone Laboratories.

SCHNETTLER, F. J., see Van Uitert, L. G.

SEIDEL, H.¹

Ferrite Slabs in Transverse Electric Mode Wave Guide, J. Appl. Phys., **28**, pp. 218-226, Feb., 1957.

SLICHTER, W. P., see tabulation at the end.

SMITS, F. M.,¹ and MILLER, R. C.¹

Rate Limitation at the Surface for Impurity Diffusion in Semiconductors, Phys. Rev., **104**, pp. 1242-1245, Dec. 1, 1956.

SUHL, H.¹

The Theory of Ferromagnetic Resonance at High Signal Powers, J. Chem. and Phys. of Solids, **1**, pp. 209-227, Jan., 1957.

SUHL, H., see Clogston, A. M.

TIEN, P. K.¹

The Backward-Travelling Power in the High Power Travelling-Wave Amplifiers, Proc. I.R.E., Letter to the Editor, **45**, p. 87, Jan., 1957.

THOMAS, D. G.,¹ and LANDER, J. J.¹

Hydrogen as a Donor in Zinc Oxide, J. Chem. Phys., **25**, pp. 1136-1142, Dec., 1956.

THURMOND, C. D.,¹ TRUMBORE, F. A.,¹ and KOWALCHIK, M.¹

Germanium Solidus Curves, J. Chem. Phys., Letter to the Editor, **25**, pp. 799-800, Oct., 1956.

THURSTON, R. N.¹

Disk-Loaded Torsional Wave Delay Line. II — Theoretical Interpretation of Results and Design Information, J. Acous. Soc. Am., **29**, pp. 20-25, Jan., 1957.

THURSTON, R. N., see Andreatch, P., Jr.

¹ Bell Telephone Laboratories.

TREUTING, R. G., see Prince, E.

TRUMBORE, F. A., see Thurmond, C. D.

VAN UITERT, L. G.,¹ READ, M. H.,¹ and SCHNETTLER, F. J.¹

Permanent Magnet Oxides Containing Divalent Metal Ions—I, J.
Appl. Phys., Letter to the Editor, **28**, pp. 280–281, Feb., 1957.

VAN UITERT, L. G.,¹ see tabulation at the end.

WALKER, L. R., see Clogston, A. M.

WENNY, D. H., see Gould, H. L.

WERNICK, J. H., see tabulation at the end.

WERTHEIM, G. K.,¹ and AUGUSTYNIAK, W. M.¹

Measurement of Short Carrier Lifetimes, Rev. Sci. Instr., **27**, pp. 1062–1064, Dec., 1956.

WOOD, E. A., see Matthias, B. T.

YERKES, E. P.⁷

Comfort A. Adams 1956 Edison Medalist — Medal History, Elec. Engg., **76**, p. 224, March, 1957.

THE ENCYCLOPEDIA OF CHEMISTRY — BOOK PUBLISHED BY REINHOLD PRESS, NEW YORK, JANUARY, 1957.

BRIDGERS, H. E., **Germanium and Its Compounds**, p. 445, and **Semiconductors**, pp. 852–583.

EDELSON, D., **Polar Molecules**, pp. 767–769.

FLASCHEN, S. S., **Calcination**, pp. 161–162.

GARN, P. D., **Electrolysis**, pp. 341–342.

HAWKINS, W. L., **Autoxidation**, p. 116.

¹ Bell Telephone Laboratories.

⁷ Bell Telephone Company of Pennsylvania, Philadelphia.

KUNZLER, J. E., **Calorimetry**, pp. 163-165.

LAW, J. T., **Vacuum Techniques**, pp. 964-965.

LUNDBERG, C. C., **Antiozonants**, pp. 97-98, and **Solutions**, pp. 873-874.

NELSON, L. S., **Olefin Compounds**, pp. 661-663.

POTTER, J. F., **Porosity**, pp. 775-776.

REISS, H., **Thermodynamics**, pp. 930-933.

SCHLABACH, T. D., **Photometric Analysis**, pp. 735-736.

SLICHTER, W. P., **Paramagnetism**, p. 700.

VAN UITERT, L. G., **Equilibrium**, pp. 363-364.

WERNICK, J. H., **Carbonates**, pp. 173-174.

Recent Monographs of Bell System Technical Papers Not Published in This Journal*

ALLISON, H. W., see Moore, G. E.

ARNOLD, S. M.

Growth and Properties of Metal Whiskers, Monograph 2635.

AUGUSTYNIAK, W. M., see Wertheim, G. K.

BASHKOW, T. R.

Effect of Nonlinear Collector Capacitance on Rise Time, Monograph 2742.

BÖMMEL, H. E., see Mason W. P.

BOZORTH, R. M.

Ferromagnetism, Monograph 2679.

BRATTAIN, W. H.

Development of Concepts in Semiconductor Research, Monograph 2743.

BRIDGERS, H. E., and KOLB, E. D.

Distribution Coefficient of Boron in Germanium, Monograph 2684.

CHEN, W. H., see Lee, C. Y.

COMPTON, K. G.

Potential Criteria for Cathodic Protection of Lead Cable Sheath, Monograph 2655.

* Copies of these monographs may be obtained on request to the Publication Department, Bell Telephone Laboratories, Inc., 463 West Street, New York 14, N. Y. The numbers of the monographs should be given in all requests.

DAVID, E. E., and McDONALD, H. S.

Note on Pitch-Synchronous Processing of Speech, Monograph 2744.

DEHN, J. W., and HERSEY, R. E.

Recent New Features for No. 5 Crossbar Switching System, Monograph 2745.

FAY, C. E.

Ferrite-Tuned Resonant Cavities, Monograph 2713.

FRY, T. C.

The Automatic Computer in Industry, Monograph 2755.

GILBERT, E. N.

Enumeration of Labelled Graphs, Monograph 2680.

GOLDEY, J. M., see Moll, J. L.

HAGSTRUM, H. D.

Auger Ejection of Electrons from Molybdenum by Noble Gas Ions, Monograph 2716.

HAGSTRUM, H. D.

Metastable Ions of the Noble Gases, Monograph 2714.

HERSEY, R. E., see Dehn, J. W.

HOLONYAK, N., see Moll, J. L.

JAYCOX, E. K., and PRESCOTT, B. E.

Spectrochemical Analysis of Thermionic Cathode Nickel Alloys, Monograph 2756.

KOLB, E. D., see Bridgers, H. E.

KRUSEMEYER, H. J., and PURSLEY, M. V.

Donor Changes in Oxide-Coated Cathodes, Monograph 2717.

LANDER, J. J., see Thomas, D. G.

LEE, C. Y., and CHEN, W. H.

Several-Valued Combinational Switching Circuits, Monograph 2746.

LOVELL, L. C., see Vogel, F. L., Jr.

MASON, W. P.

Internal Friction and Fatigue in Metals at Large Strain Amplitudes, Monograph 2758.

MASON, W. P.

Physical Acoustics and Properties of Solids, Monograph 2761.

MASON, W. P., and BÖMMEL, H. E.

Ultrasonic Attenuation at Low Temperatures for Metals, Normal and Superconducting, Monograph 2748.

MATLACK, R. C.

Role of Communications Networks in Digital Data Systems, Monograph 2678.

METREYEK, W., see Winslow, F. H.

MCDONALD, H. S., see David, E. E.

McSKIMIN, H. J.

Wave Propagation and Measurement of Elasticity of Liquids and Solids, Monograph 2749.

MERZ, W. J.

Switching Time in Ferroelectric BaTiO₃ and Crystal Thickness, Monograph 2721.

MILLER, R. C., and SAVAGE, A.

Diffusion of Aluminum in Single-Crystal Silicon, Monograph 2722.

MOLL, J. L., TANENBAUM, M., GOLDEY, J. M., and HOLONYAK, N.

P-N-P-N Transistor Switches, Monograph 2723.

MOORE, G. E., and ALLISON, H. W.

Emission of Oxide Cathodes Supported on a Ceramic, Monograph 2724.

MOSHMAN, J., see Tien, P. K.

OHM, E. A.

A Broad-Band Microwave Circulator, Monograph 2726.

OWENS, C. D.

Modern Magnetic Ferrites and Their Engineering Applications, Monograph 2709.

OWENS, C. D.

Properties and Applications of Ferrites Below Microwave Frequencies, Monograph 2727.

PATERSON, E. G. D.

Nike Quality Assurance, Monograph 2728.

PIERCE, J. R., and WALKER, L. R.

Growing Electric Space-Charge Waves, Monograph 2729.

PIERCE, J. R.

Instability of Hollow Beams, Monograph 2751.

PRESCOTT, B. E., see Jaycox, E. K.

PURSLEY, M. V., see Krusemeyer, H. J.

ROSE, D. J.

Townsend Ionization Coefficient for Hydrogen and Deuterium, Monograph 2731.

SAVAGE, A., see Miller, R. C.

SLEPIAN, D.

Note on Two Binary Signaling Alphabets, Monograph 2733.

SPROUL, P. T.

A Video Visual Measuring Set with Sync Pulses, Monograph 2752.

TANENBAUM, M., see Moll, J. L.

THOMAS, D. G., and LANDER, J. J.

Hydrogen as a Donor in Zinc Oxide, Monograph 2753.

TIEN, P. K., and MOSHMAN, J.

Noise in a High-Frequency Diode, Monograph 2735.

TORREY, M. N.

Quality Control in Electronics, Monograph 2736.

VAN UITERT, L. G.

Dielectric Properties of and Conductivity in Ferrites, Monograph 2737.

VOGEL, F. L., JR., and LOVELL, L. C.

Dislocation Etch Pits in Silicon Crystals, Monograph 2738.

WALKER, L. R., see Pierce, J. R.

WEINREICH, G.

Acoustodynamic Effects in Semiconductors, Monograph 2764.

WEISS, M. T.

Improved Rectangular Waveguide Resonance Isolators, Monograph 2739.

WERTHEIM, G. K.

Carrier Lifetime in Indium Antimonide, Monograph 2740.

WERTHEIM, G. K., and AUGUSTYNIAK

Measurement of Short Carrier Lifetimes, Monograph 2754.

WINSLOW, F. H., and MATREYEK, W.

Pyrolysis of Crosslinked Styrene Polymers, Monograph 2741.

Contributors to This Issue

KENNETH BULLINGTON, B.S., University of New Mexico, 1936; M.S., Massachusetts Institute of Technology, 1937; Bell Telephone Laboratories, 1937-. Mr. Bullington's first work with the Laboratories was on systems engineering on wire transmission circuits, and since 1942 he has been concerned with transmission engineering on radio systems, particularly over-the-horizon radio propagation. In 1956, he was awarded the Morris Liebmann Memorial Prize of the I.R.E. and the Stuart Ballentine Medal from the Franklin Institute for contributions in tropospheric transmission and the application of those contributions to practical communication systems. He is a Fellow of the I.R.E., and a member of Phi Kappa Phi, Sigma Tau and Kappa Mu Epsilon.

ARTHUR B. CRAWFORD, B.S.E.E. 1928, Ohio State University; Bell Telephone Laboratories 1928-. Mr. Crawford has been engaged in radio research since he joined the Laboratories. He has worked on ultra short wave apparatus, measuring techniques and propagation; microwave apparatus, measuring techniques and radar; microwave propagation studies and microwave antenna research. He is author or co-author of articles which appeared in the Bell System Technical Journal, Proceedings of the I.R.E., Nature and Bulletin of the American Meteorological Society. He is a Fellow of the I.R.E. and a member of Sigma Xi, Tau Beta Pi, Eta Kappa Nu, and Pi Mu Epsilon.

HAROLD E. CURTIS, B.S. and M.S., Massachusetts Institute of Technology, 1929; Department of Development and Research of the American Telephone and Telegraph Company, 1929; Bell Telephone Laboratories, 1934-. Mr. Curtis has been concerned with transmission problems related to multi-channel carrier telephony. He has also been engaged in studies of transmission engineering aspects of the microwave radio relay system. His work at the Laboratories has also included pioneering transmission studies of the coaxial cable, the shielded pair and quad, and the waveguide. Mr. Curtis holds ten patents relating to carrier telephony.

HARALD T. FRIIS, E.E., 1916, D.Sc., 1938, Royal Technical College (Copenhagen); Western Electric Company, 1919; Bell Telephone Lab-

oratories, 1930-. Dr. Friis, Director of Research in High Frequency and Electronics, has made important contributions on ship-to-shore radio reception, short-wave studies, radio transmission (including methods of measuring signals and noise), a receiving system for reducing selective fading and noise interference, microwave receivers and measuring equipment, and radar equipment. He has published numerous technical papers and is co-author of a book on the theory and practice of antennas. The I.R.E.'s Morris Liebmann Memorial Prize, 1939, and Medal of Honor, 1954. Valdemar Poulson Gold Medal by Danish Academy of Technical Sciences, 1954. Danish "Knight of the Order of Dannebrog," 1954. Fellow of I.R.E. and A.I.E.E. Member of American Association for the Advancement of Science, Danish Engineering Society and Danish Academy of Technical Sciences. Served on Panel for Basic Research of Research and Development Board, 1947-49, and Scientific Advisory Board of Army Air Force, 1946-47.

LESTER H. GERMER, B.A., Cornell, 1917; M.A., Columbia, 1922; Ph.D., Columbia, 1927; Western Electric Co., 1917-24; Bell Telephone Laboratories, 1925-. With the Research Department, Dr. Germer has been concerned with studies in electron diffraction, structure of surface films, thermionics, contact physics, order-disorder phenomena, and physics of arc formation. He has published about seventy papers and has three patents. In 1931 he received the Elliott Cresson medal of the Franklin Institute. He is a member of the American Physical Society, Sigma Xi, the New York Academy of Sciences, the A.A.A.S., and the American Crystallographic Society of which he served as president in 1944.

DAVID C. HOGG, B.S., University of Western Ontario, 1949; M.S. and Ph.D., McGill University, 1950 and 1953; Bell Telephone Laboratories, 1953-. Mr. Hogg has been engaged in studies of artificial dielectrics for microwaves, antenna problems, and over-the-horizon and millimeter wave propagation as a member of the Radio Research Dept. During World War II, Mr. Hogg served with the Canadian Army in Europe and from 1950-51 did research for the Defense Research Board of Canada. He is a member of Sigma Xi, and a senior member of the I.R.E.

STEPHEN O. RICE, B.S., Oregon State College, 1929; California Institute of Technology, Graduate Studies, 1929-30 and 1934-35; Bell Telephone Laboratories, 1930-. In his first years at the Laboratories, Mr. Rice was concerned with the non-linear circuit theory, with special

emphasis on methods of computing modulation products. Since 1935 he has served as a consultant on mathematical problems and in investigations of the telephone transmission theory, including noise theory, and applications of electromagnetic theory. Fellow of the I.R.E.

J. W. SCHAEFER, B.M.E., Ohio State University, 1941; Bell Telephone Laboratories, 1940-. Mr. Schaefer has worked on dial design and dial test equipment, and during the war years contributed to the design and development of anti-aircraft fire control equipment and guided missiles. After the war, Mr. Schaefer proposed a means of steering missiles from which evolved NIKE. He is now working on anti-aircraft guided missile systems. He is a member of A.S.M.E., the Army Ordnance Association, Tau Beta Pi and Sigma Xi.

BERNARD SMITH, B.S., City College of New York, 1948; A.M., 1951, and Ph.D., 1954, Columbia University; Lecturer, City College of New York, 1948-1954; Bell Telephone Laboratories, 1954-. In addition to the transmission studies in which he has been engaged since joining the Laboratories, his present duties include teaching information theory in the Communications Development Training Program. He is a member of the American Physical Society, Phi Beta Kappa, Sigma Xi and Kappa Delta Pi.

JAMES L. SMITH, B.S., Newark College of Engineering, 1956; Bell Telephone Laboratories, 1941-. Mr. Smith worked on problems concerned with relay contact erosion as a technical aide, and in 1956 began his work on solid state switching networks. He is a member of the A.I.E.E. and Tau Beta Pi.

MARK A. TOWNSEND, B.S., Texas Technological College, 1936; M.S., Mass. Institute of Technology, 1937; Bell Telephone Laboratories, 1945-. Mr. Townsend's early work with the Laboratories was on the development of gas discharge tubes for use in telephone switching systems. More recently, his work has been in the exploratory development of systems for digital data transmission and of a small electronic switching system. He is a member of the A.I.E.E., and senior member of the I.R.E.

GERD F. WEISSMANN. Dipl.-Ing. Technical University of Berlin, 1950; M.S. Pennsylvania State University, 1953; Bell Telephone Laboratories, 1953-. Mr. Weissmann's work at the Laboratories has been in stress analysis, engineering mechanics, strain measurements, soil mechanics and metals properties and testing. He also has worked with outside plant problems and metallurgical engineering.

H E B E L L S Y S T E M

Technical Journal

DEVOTED TO THE SCIENTIFIC AND ENGINEERING
ASPECTS OF ELECTRICAL COMMUNICATION

VOLUME XXXVI

JULY 1957

NUMBER 4

Noise Spectrum of Electron Beam in Longitudinal
Magnetic Field W. W. RIGROD

Part I—The Growing Noise Phenomenon

831

Part II—The UHF Noise Spectrum

855

Distortion Produced in a Noise Modulated FM Signal by Non-
linear Attenuation and Phase Shift

S. O. RICE 879

Self-Timing Regenerative Repeaters

E. D. SUNDE 891

A Sufficient Set of Statistics for a Simple Telephone Exchange
Model

V. E. BENEŠ 939

Fluctuations of Telephone Traffic

V. E. BENEŠ 965

High-Voltage Conductivity-Modulated Silicon Rectifier

H. S. VELORIC AND M. B. PRINCE 975

Coincidences in Poisson Patterns E. N. GILBERT AND H. O. POLLAK 1005

Bell System Technical Papers Not Published in This Journal 1035

Recent Bell System Monographs 1043

Contributors to This Issue 1045

THE BELL SYSTEM TECHNICAL JOURNAL

ADVISORY BOARD

A. B. GOETZE, *President, Western Electric Company*

M. J. KELLY, *President, Bell Telephone Laboratories*

E. J. MCNEELY, *Executive Vice President, American Telephone and Telegraph Company*

EDITORIAL COMMITTEE

B. MCMILLAN, *Chairman*

S. E. BRILLHART

A. J. BUSCH

L. R. COOK

A. C. DICKIESON

R. L. DIETZOLD

K. E. GOULD

E. I. GREEN

R. K. HONAMAN

H. R. HUNTLEY

F. R. LACK

J. R. PIERCE

G. N. THAYER

EDITORIAL STAFF

W. D. BULLOCH, *Editor*

R. L. SHEPHERD, *Production Editor*

T. N. POPE, *Circulation Manager*

THE BELL SYSTEM TECHNICAL JOURNAL is published six times a year by the American Telephone and Telegraph Company, 195 Broadway, New York 7, N. Y. F. R. Kappel, President; S. Whitney Landon, Secretary; John J. Scanlon, Treasurer. Subscriptions are accepted at \$5.00 per year. Single copies \$1.25 each. Foreign postage is 65 cents per year or 11 cents per copy. Printed in U. S. A.

THE BELL SYSTEM TECHNICAL JOURNAL

VOLUME XXXVI

JULY 1957

NUMBER 4

Copyright 1957, American Telephone and Telegraph Company

Noise Spectrum of Electron Beam in Longitudinal Magnetic Field

By W. W. Rigrod

(Manuscript received January 21, 1957)

Measurements of induced noise currents along drifting cylindrical electron beams have shown that noise fluctuations propagate as space-charge waves in the same fashion as RF signals of the same frequency. On many such beams, however, the regular standing-wave noise pattern is interrupted, after some drift distance, by a smooth steep increase in noise current, followed by slow, shallow undulations. This "growing noise" phenomenon, discovered by Smullin and his co-workers at M.I.T. several years ago, is the subject of study in this paper. Its importance is considerable, in a negative way, because it has hampered the development of medium-power traveling-wave-tube devices with acceptably low noise figures.

The experimental measurements show the growing noise pattern to be the result of a two-stage process. Its primary cause is rippled-beam amplification of noise fluctuations over a wide band of microwave frequencies, much higher than the usual observation frequency. This explains its elusiveness. In the second stage, noise energy is transferred to lower frequencies, due to intermodulation and other non-linear processes within the gain band. As the beat-frequency noise increments are excited by continuous arrays of frequency pairs, their standing-wave patterns overlap one another, resulting in a smooth growing-noise pattern.

In Part II of this paper, measurements of the noise spectrum of a rippled beam in the UHF region are described. These measurements reveal the presence of additional forms of instability. Calculations are made to account for some of these, and for aspects of rippled-beam amplification not previously understood.

Part I — The Growing Noise Phenomenon*

I INTRODUCTION

When an RF probe is moved along a magnetically-focused electron beam in a drift region, the noise power is at first found to vary periodically with distance from the electron gun.¹ For a sufficiently long beam, however, the periodic pattern is succeeded by an exponential rise, culminating in an irregular plateau. This so-called "growing noise" phenomenon has been extensively investigated by its discoverers, L. Smullin and his colleagues at the M.I.T. Research Laboratory of Electronics.^{2, 3} They have established that this noise will begin to grow at a plane nearer the gun, and tend to grow at a faster rate, for electron beams (a) of higher perveance, (b) with less space-charge neutralization by positive ions, and (c) issuing from convergent, partly-shielded guns, rather than those immersed in the magnetic field.

The growth of microwave noise power in drifting beams has hampered the development of high-power, traveling-wave tubes with acceptably low noise figures, as such devices generally have convergent, partly-shielded electron guns. The problem has been evaded in the design of low-noise, low-power traveling-wave tubes, by resort to confined-flow, parallel beams.

Several theories have been proposed to explain the growing-noise wave:

- (1) Excitation of higher-order modes with complex propagation constants, by electrons threading the beam transversely;⁴
- (2) Slipping-stream amplification, due to either longitudinal or transverse velocity gradients;⁵
- (3) Rippled-beam amplification;^{6, 7, 8} and
- (4) Electron-electron interactions leading eventually to equipartition of thermal energy, and thus an increase in longitudinal velocity fluctuations.

In Part I of this paper, measurements are presented which show that the principal cause of growing noise appears to be space-charge wave

* Presented at the I.R.E. Electron Tube Research Conference, Boulder, Colorado, June 27-29, 1956.

amplification due to beam rippling. The mechanism is studied in some detail, as its connection with the usually-observed exponential rise of noise is not immediately apparent. In Part II, the UHF noise spectrum and its spatial distribution in beams with large-amplitude, long wavelength ripples, are described. In addition, some of the underlying processes are analyzed.

II APPARATUS

As sketched in Fig. 1, the heart of the apparatus consists of an electron gun, drift tube, and movable probe, all enclosed in a demountable, continuously-pumped vacuum system. Outside of the vacuum envelope there is a shielded solenoid, extending the entire 18-inch length of the drift tube. The annular gap between the solenoid pole face and the magnetic shield about the gun is nearly all taken up by a soft-steel section of the vacuum envelope.

The electron gun is of the convergent Pierce type, with oxide-coated cathode and a coiled-coil filament heater producing negligible flux at the cathode surface. Surrounding the gun, and inside of the magnetic shield, is a small copper-wire coil that permits variation of this flux over a small range, either aiding or opposing the leakage flux due to the main focusing solenoid. The flux density at the cathode has been approximately calibrated in terms of currents in both coils. Throughout the experiments described below, the gun is pulsed with a 1,000 cps square wave of 2,200 volts on its anode, supplying 38 ± 1 ma peak current in space-charge-limited emission.

The novel feature of the probe is that its annular RF pickup gap couples to a 50-ohm coaxial line leading to the receiver, rather than to a resonant cavity. This permits RF power measurements over a wide range of frequencies. The inner conductor of the coaxial line serves as current-collector, being isolated and biased positively about 40 volts with respect to the outer conductor to prevent escape of slow secondaries. An adjustable vane can be locked in position in front of the probe (whose entrance aperture is 0.100 inch in diameter), so that circular apertures of various smaller sizes are fixed on the probe centerline, about 0.070 inch in front of the probe. With these apertures, measurements of collector-current variations along the beam furnish a rough picture of beam-ripple amplitudes and locations. In addition, the current-density variation across the beam can be estimated by moving a pinhole aperture in a broad arc through the beam centerline. Both the inner conductor of the probe and the intercepting vane are liquid-cooled.

The noise powers coupled to the coaxial probe are considerably smaller

than for a tuned coaxial cavity, because of the lower RF gap impedance of the former. To compensate for this drawback, a sensitive noise receiver is employed, similar in principle to the radiometer invented by R. H. Dicke.⁹ The input noise power is replaced periodically by a matched load at room temperature by pulsing the beam on and off with a 1,000 cps square wave, and placing an isolator in front of the receiver. A synchronous detector eliminates gain-fluctuation noise and converts the receiver output to a dc voltage.

Noise power variations at various microwave frequencies are measured in terms of the changes of attenuation, between probe and receiver, required to keep the receiver output constant. These rapid adjustments in attenuation are performed by a servo amplifier-motor loop, and recorded on a chart, whose speed ($1\frac{1}{2}$ inches per minute) is synchronized with that of the moving probe. In the same way, records of collector current as a function of probe position can be obtained, and correlated with those of noise power. The probe can be moved a distance of about 17 inches, its position nearest the gun ($z = 0$ inches) corresponding to a distance of 0.95 inch between the anode and the input plane of the RF gap.

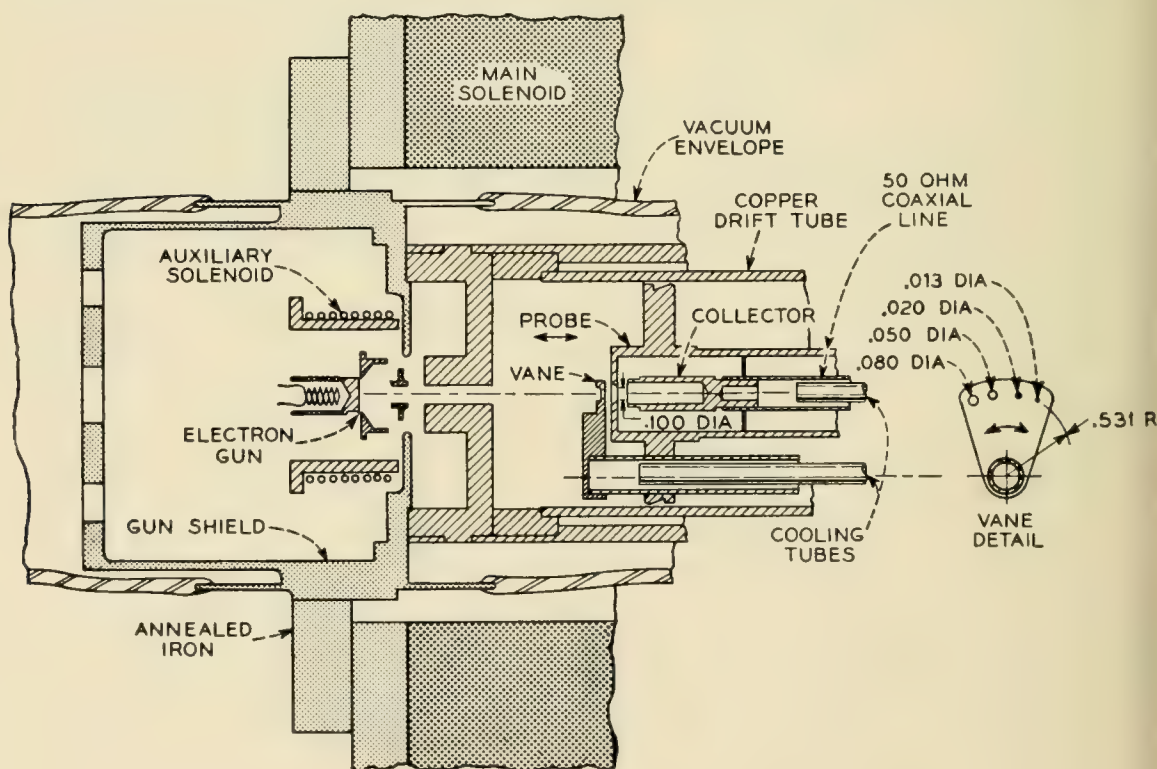


FIG. 1 — Cross-section of experimental tube, showing electron gun, probe, and two solenoids. The isolated current-collector electrode serves as inner conductor of a coaxial line. The induced RF power can be measured over a wide range of frequencies.

III EXPLORATORY MEASUREMENTS

The electron gun used in these experiments had been designed for use in a helix traveling-wave tube with a longitudinal focusing field of 600 gauss. Noise-power and collector-current curves, therefore, were first taken with 600 gauss to study a typical state of affairs in an operational beam. As seen in Fig. 2, the noise power at 3.9 kmc varies periodically with distance for about 4 inches from the gun, then climbs rather smoothly by nearly 23 db to an irregular plateau, where it undulates slowly, and finally levels off. The initial part of the growing noise curve at 10.7 kmc is missing because of inadequate receiver sensitivity, but its later portion is similar to that at 3.9 kmc, with about half the rate of noise climb. With the 0.020-inch aperture, the collector-current variations decrease in amplitude chiefly in the drift region preceding the noise climb; whereas those for the 0.100-inch aperture decrease afterwards. Both curves show a flattening in the growing-noise region itself, as well as a decrease in their *average* values after that region, signifying an increase in the average beam diameter.

A similar set of curves is shown in Fig. 3, for a focusing field of 279 gauss (about twice the nominal Brillouin field). Noise growth at 3.9 kmc starts later, and proceeds less steeply, than at 600 gauss. The noise-power curve for 10.7 kmc is much more articulated, with a semblance of periodicity, throughout the drift region. Collector-current curves for both 0.020- and 0.050-inch apertures show considerable reduction in current-ripple amplitude with distance, reaching virtually zero in the former case.

Another type of survey measurement is illustrated in Fig. 4. With the probe stationary at the far end of the drift space (about 18 inches from the gun anode), the main solenoid current is varied smoothly to change the focusing field from 0 to over 600 gauss, and synchronized records are made of collector current and noise power. (In this instance, the current in the auxiliary solenoid was +3.2 amperes.) At low magnetic fields, both the current and noise-power curves have large amplitude variations, which diminish as the field increase. At first glance, the noise peaks and valleys seem to coincide with those of collector current; certainly, some do. Closer inspection, however, reveals significant misalignments which cannot be accounted for by experimental error. When the three noise curves, at 3,050, 3,930, and 4,730 mc, respectively, are compared with each other, some characteristic features emerge:

(1) An average curve drawn through each pattern has one or two broad maxima, which tend to move toward higher field strengths with increasing frequency.

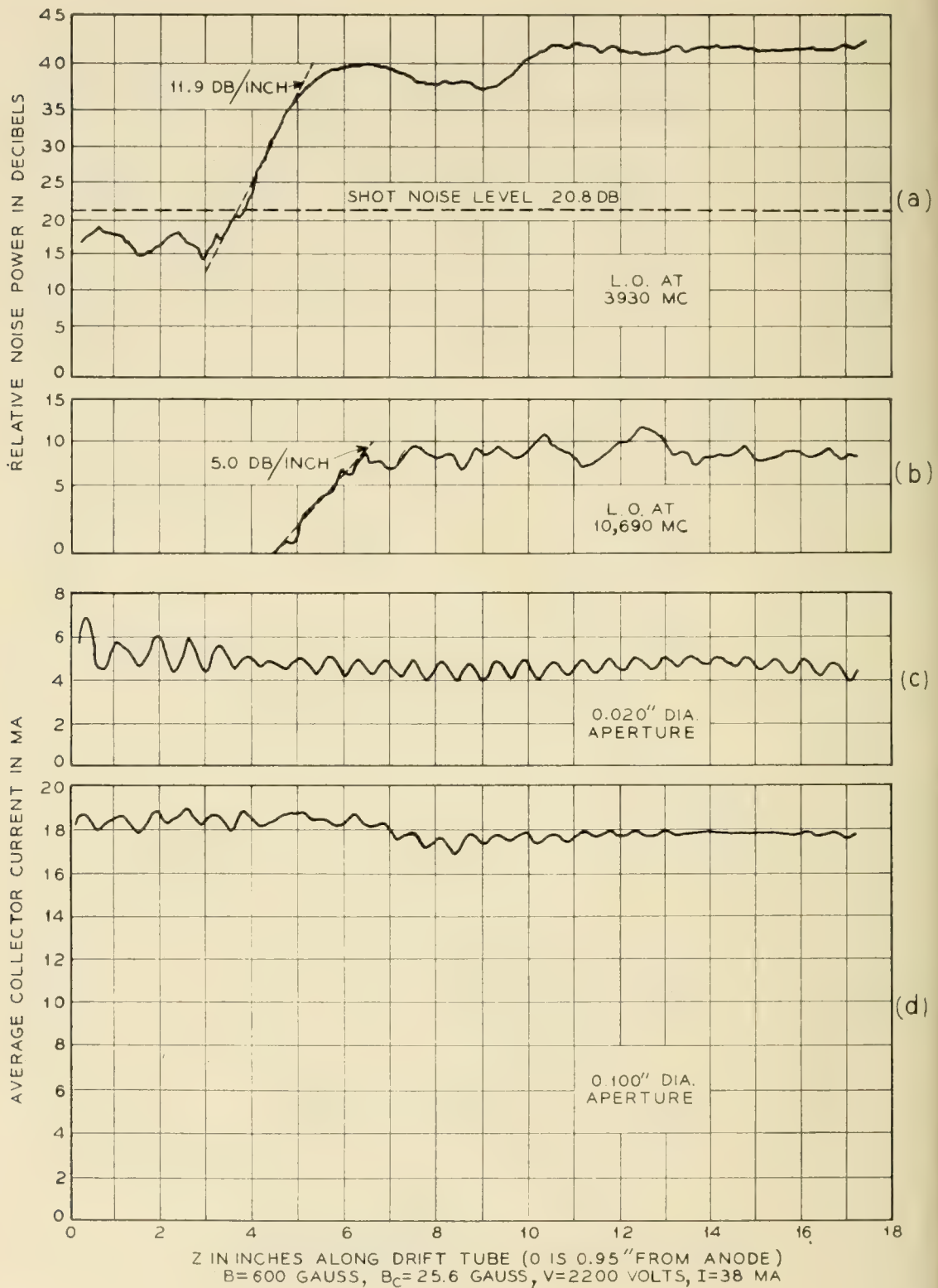


FIG. 2 — Typical smooth steep growing-noise patterns, near 4 and 10.7 kmc, respectively, with customary focusing field of about four times the Brillouin value. Collector-current traces through small and large apertures reveal decreases in ripple amplitude and increase in average beam diameter.

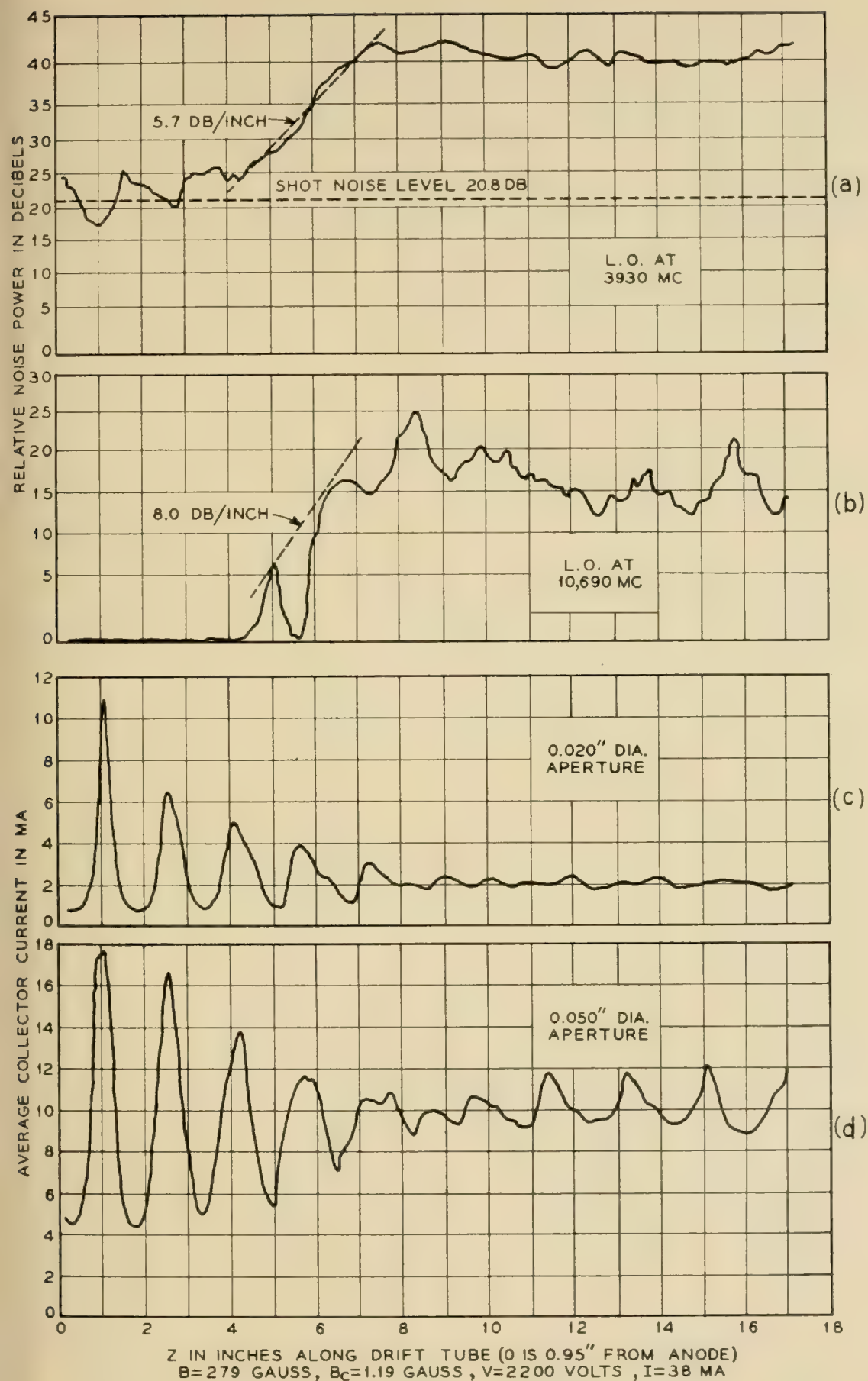


FIG. 3 — At about half the focusing fields used in Fig. 2, the growing-noise pattern is much the same at 4 kmc, but shows significant articulation at 10.7 kmc. The collector-current traces show pronounced decreases in ripple amplitude, and differences in ripple patterns obtained with different apertures.

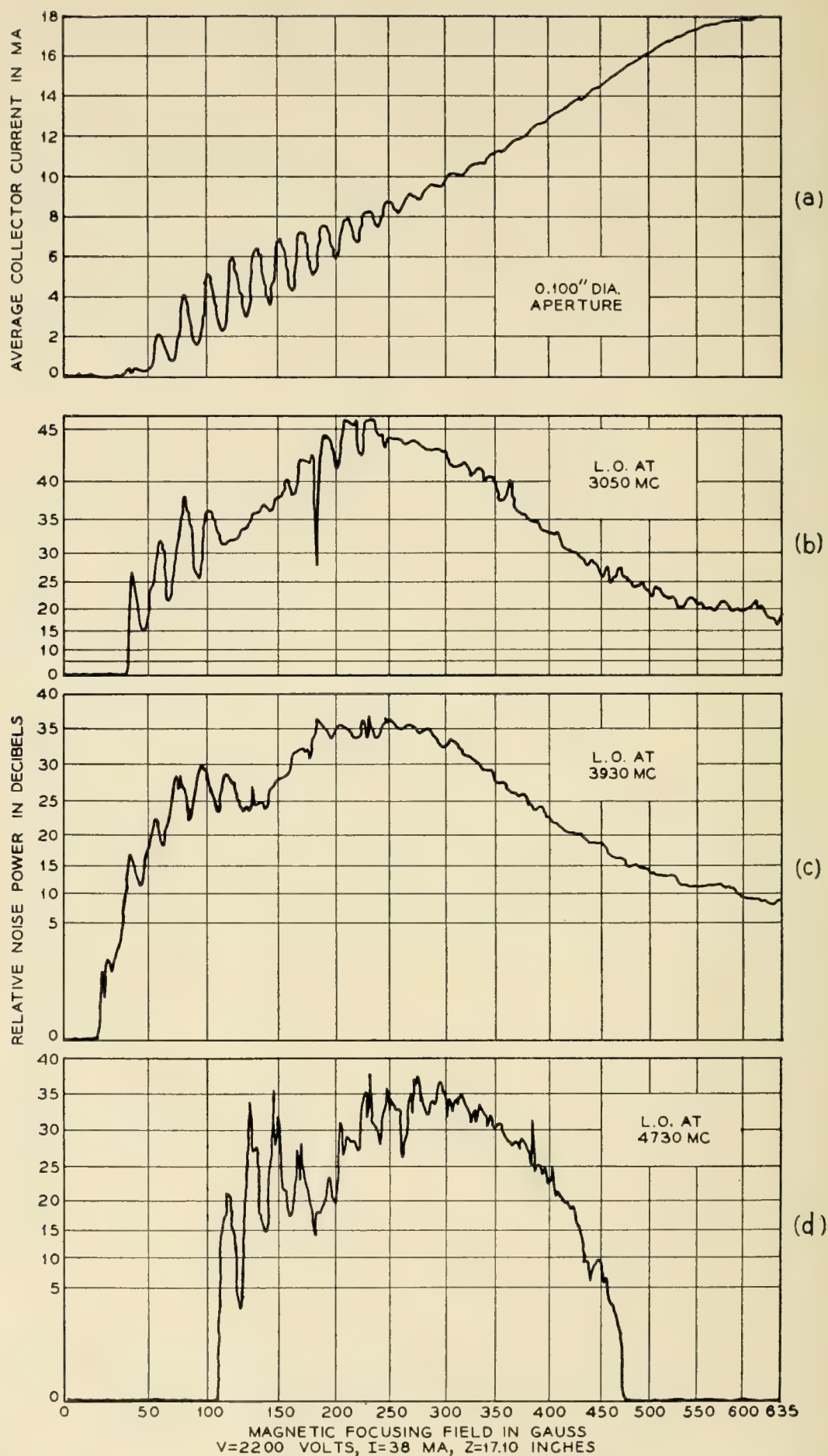


FIG. 4 — With the probe stationary at about 18 inches from gun anode, collector current and noise powers in bands about 1 kmc apart are recorded, as current in the main solenoid is varied from 0 to 1.5 amperes (0 to 635 gauss).

(2) The lower the frequency, the lower the field strengths at which noise amplitudes change most violently with field.

(3) The three noise curves resemble each other in small details.

The results of a great many records of the kind illustrated by Figs. 2 and 3 can be summarized as follows:

(1) There is always a decrease in beam-ripple amplitude associated with noise growth at any frequency. (Sometimes the ripple amplitude increases afterwards, as in Fig. 3.)

(2) The higher the frequency, for a given field, the more articulated or scalloped the noise pattern.

(3) No correlation can be found between rate of noise growth and either (a) distance from gun to take-off plane, or (b) net gain at the end of the drift region. The trends, as a function of magnetic field, are different at different frequencies.

(4) Greatest noise growth does not, as a rule, occur with zero flux threading the cathode. Sometimes two nearly equal peaks occur for two values of B_c , each of opposite polarity, referred to the sense of the main field.

(5) The noise-distance patterns change very slowly with frequency.

(6) No beam entirely ripple-free throughout its length has ever been observed by the writer.

IV ORIGIN OF GROWING NOISE

If noise growth is due to some amplification process, it should be possible to adjust the beam-focusing conditions so that the noise currents start increasing at the anode, and attain the greatest possible over-all gain at the end of the drift space. The enhanced activity of the unknown gain mechanism should presumably help identify it. The curves of Fig. 4 show that maximum noise occurs at different values of the focusing field, for different values of field at the cathode, and different probe positions. With the anode voltage and receiver frequency fixed, therefore, the conditions for greatest net noise growth can only be found by a series of trial settings of *both* magnetic fields, each followed by a recording of the noise-distance pattern. Eventually, a set of fields can be found for which the greatest total gain occurs; and such patterns are usually found to show fairly steady noise-amplitude increase, on the average, over the entire length of probe travel.

The results of this procedure for noise power near 4 kmc, as well as the patterns of collector-current versus distance with the same fields, using the 0.100-, 0.050-, and 0.020-inch apertures fixed at the probe centerline, are shown in Fig. 5. A similar set of records, for noise power

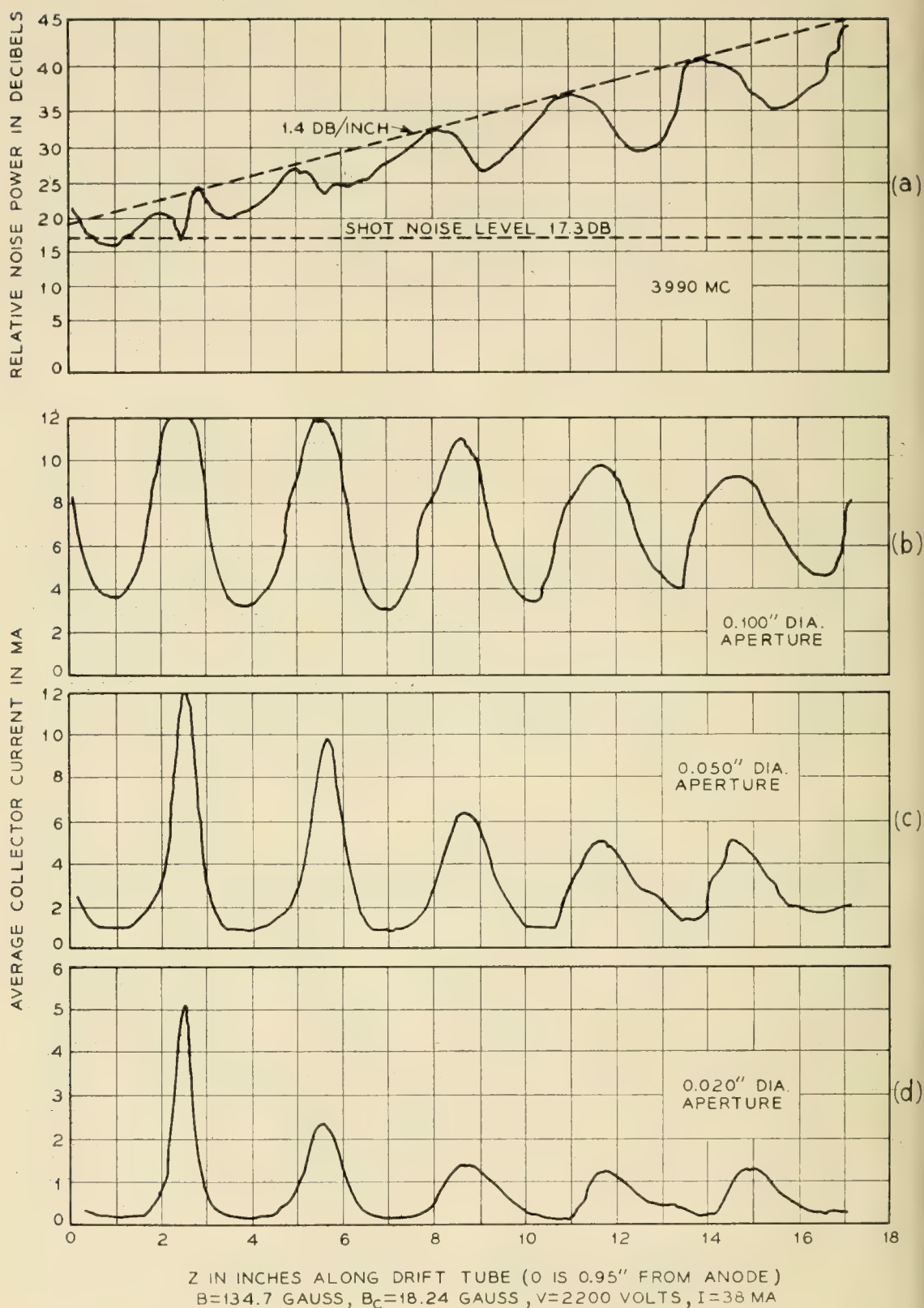


FIG. 5 — The magnetic fields in drift tube and at the cathode have been adjusted empirically to expand the growing-noise region over the entire drift region, with the L.O. at 3,990 mc. The field is slightly less than the Brillouin value, but the beam is strongly rippled because the gun was designed for best focusing at a much higher field. The noise-current maxima align with the average collector currents on their increasing slopes, for all three aperture sizes.

near 10.69 kmc, is shown in Fig. 6. The significant features of both sets of records can be summarized as follows:

(1) In both cases, the beam-ripple periods are equal to the RF scallop periods; i.e., the half-wavelengths of the space-charge standing waves. The noise minima tend to occur at planes where the collector currents are at their average values and decreasing; i.e., where the beam diameters are at their average values and increasing. The noise-current maxima occur where the beam diameters are about to decrease. These are the classical conditions for *rippled-beam amplification*.^{6, 7, 8}

(2) In Fig. 5, the ripple amplitudes and peak values of all three collector-current curves decline appreciably with distance, the rate of decline being greatest for the smallest aperture. (Similar curves, not shown here, have displayed little or no such decline in the absence of noise growth.) This suggests that the RF noise power is amplified at the expense of dc energy associated with radially-directed electron velocities.

(3) In Fig. 6, the disparity among rates of decline of current-ripple amplitudes and their peak values, for the three aperture sizes, is even more pronounced. In addition, the ripple wavelength barely changes for the 0.100-inch aperture, but increases with drift distance for the smaller apertures, resulting in an increasing "phase shift" among them. Thus the current-density variations at different radii in the beam can contribute unequally to space-charge wave amplification, depending on their local ripple amplitude and phase. In this instance, the variations in current density along the beam are initially greatest near the axis, and suffer the greatest reduction there. It is worth noting that this "inner rippling" would be missed entirely in beam-size measurements with a large aperture.*

The decrease of beam ripple and the increase in average beam diameter, shown in Figs. 5 and 6, has been found to accompany rippled-beam amplification of impressed signals by T. G. Mihran.⁸ Another corroboration of the identity of this gain mechanism can be obtained by comparing the measured noise gain per scallop with that predicted by theory for idealized conditions.^{6, 7} For a beam with stepwise alternations of maximum and minimum beam diameters (ratio r_2/r_1), and with noise maxima and beam-diameter maxima coinciding, the gain per scallop is as follows:

$$G_m = \left(\frac{V_1}{V_2} \right)^{3/4} \left(\frac{r_2}{r_1} \right) \left(\frac{p_1}{p_2} \right). \quad (1)$$

Here, V is the beam potential, and p the reduction factor ω_q/ω_p . Although the actual rippled beam is far removed from either Brillouin or

* More information about "inner rippling" will be presented in Part II.

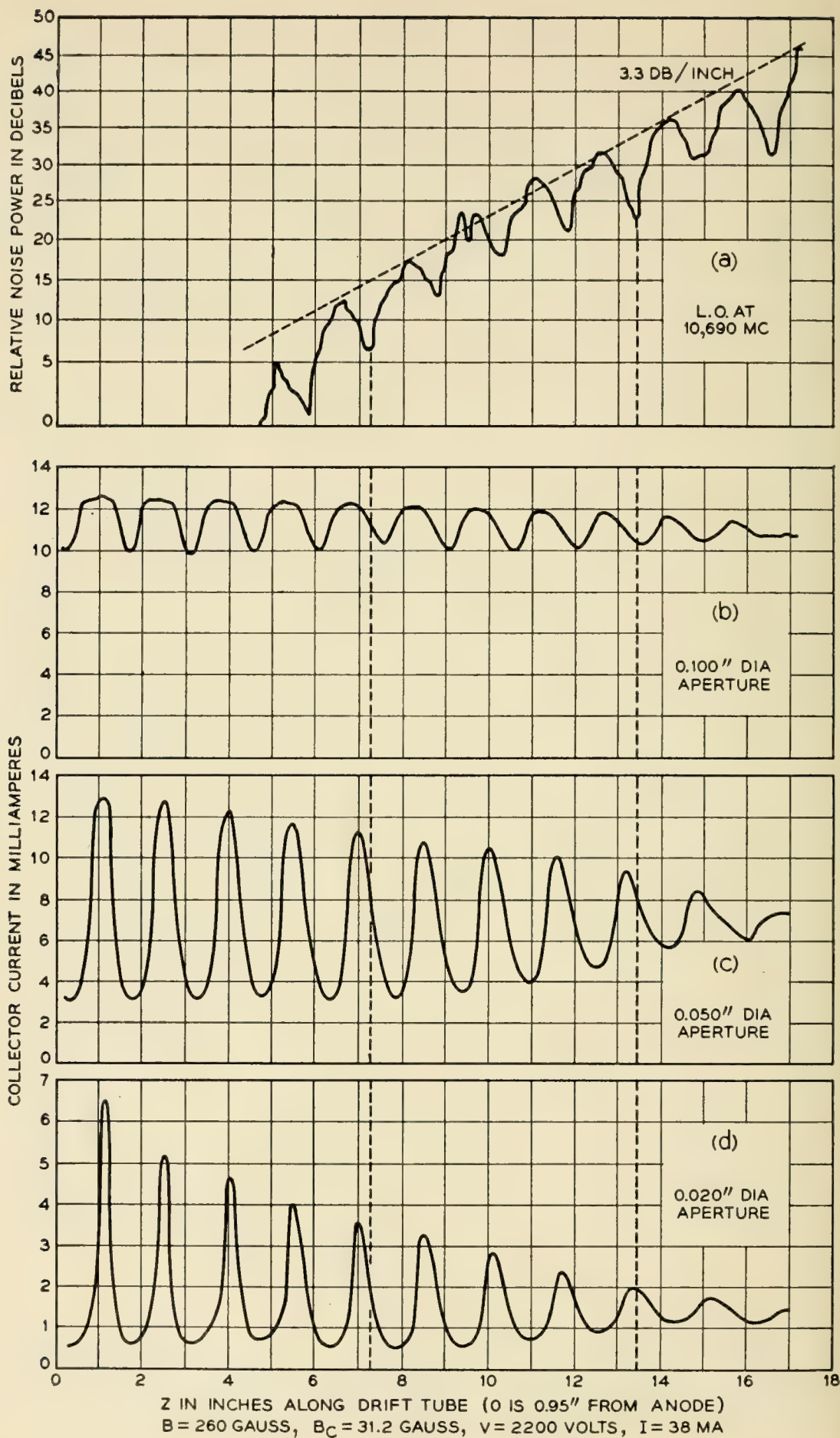


FIG. 6 — The fields have been adjusted as in Fig. 5, for maximum extension of the gain region, for noise power near 10,690 mc. Collector current measured with the smallest aperture shows the greatest decline in amplitude of variations, as well as advance in ripple phase relative to the current through the largest aperture.

TABLE I — MEASURED VERSUS CALCULATED MAXIMUM NOISE GAIN

Freq. mc.	Ripple Data		Gain in db per Scallop		
	Iris Dia. Inches	r_2/r_1	Brillouin Flow	Confined Flow	Measured
3,990	0.020	3.4	6.1	5.0	4.3
	0.050	3.2	6.3	5.4	
	0.100	1.9	4.0	3.3	
10,690	0.020	2.5	6.4	5.6	4.9
	0.050	1.8	4.8	4.4	
	0.100	1.1	<1.0	<1.0	

confined flow, the published values of reduction factor for both extremes can be used as first approximations.^{10, 11} The ratio r_2/r_1 can be estimated by assuming the current density to be uniform over the beam cross-section near the middle of the drift region, for each of the three apertures used. The potential variations can be neglected. The results of such calculations are given in Table I.

As the computed gains are expected to be somewhat greater than those measured, because of the optimum conditions assumed, the best correspondence between measured and computed gain rates appears to be for the ripple data taken with the 0.050-inch iris at 3,990 mc, and that with the 0.020-inch iris at 10,690 mc. This distinction is in accord with previous qualitative comparison of Figs. 5 and 6, showing that most of the beam cooperates in the ripples of the former, but that “inner rippling” characterizes the latter.

Another calculation that reveals which part of the beam is interacting with the RF noise field in each case is that of the space-charge half-wavelength, as follows:

$$\frac{\lambda_s}{2} \cong \frac{\pi b}{p \cdot \beta_p b}, \tag{2}$$

where
$$\beta_p b = 174 I^{1/2} / V^{3/4}. \tag{3}$$

Here β_p is the plasma wave number, b the beam radius, and p the reduction factor, which can be evaluated as previously for the smooth beam in either ideal Brillouin or confined flow. For the gun used here, the square root of the perveance is

$$I^{1/2} / V^{3/4} = 0.606 \times 10^{-3} \text{ MKS units},$$

or
$$\lambda_s/2 \cong 29.8 \, b/p. \tag{4}$$

TABLE II — MEASURED VERSUS CALCULATED SPACE-CHARGE HALF-WAVELENGTHS

Freq. mc.	Iris used, inches dia.	Avg. beam radius r_0 inches	γ rad./in.	$\lambda_s/2$, inches			Ripple Wavelength L , meas'd
				Brillouin Flow	Confined Flow	Meas'd	
3,990	0.050	0.057	22.9	3.2	2.7	3.0	3.06
10,690	0.020	0.033	61.2	1.6	1.3	1.47	1.52

Thus, agreement between this expression and the measured value requires the correct choice of the effective beam radius, b . It turns out that the suitable value for Fig. 5 (3,990 mc) is the average beam radius obtained from ripple data taken with the 0.050-inch iris, and that for Fig. 6 (10,690 mc) is obtained with data taken with the 0.020-inch iris. The results are summarized in Table II.

With this mechanism as the primary source of the noise gain, it becomes clear why nearly equal noise maxima were found, with some values of the main focusing field, B , for two values of cathode flux density B_c of opposite polarity. From approximate analyses of beam ripples when flux threads the cathode, such as those provided by McDowell¹² and others, it is found that the ripple wavelength depends nearly altogether on B . Its amplitude and spatial phase, however, depend on B_c , as this affects the beam geometry at the drift-space entrance. For a sufficiently wide range of variation of B_c , the spatial phase of the ripples can be varied from the proper relation with the space-charge standing wave for gain, through the positions for de-amplification, and back to gain again.

V THE GROWING-NOISE MECHANISM

Although many earlier noise records can be understood in the light of the rippled-beam amplification (RBA) process, this is not yet true of the smooth, steep noise growth usually observed, as in Figs. 2 and 3. The simple theory predicts that a space-charge wave will be amplified when, for small ripples, a "resonance" condition exists between the ripple wavelength, L , and the space-charge half-wavelength

$$L \cong n\lambda_s/2, \quad (5)$$

where n is an integer, usually unity. In addition, as mentioned earlier, there is an optimum phase relation between ripple and standing wave for maximum gain. These conditions are not satisfied by the records of

Figs. 2 and 3, except possibly for the 10.7 kmc noise current in Fig. 3 (at a relatively low magnetic field).

To establish a connection between the two types of noise growth, the noise record of Fig. 6 (for 10.7 kmc with greatly expanded gain region) is compared with that near 4 kmc under the same conditions, in Fig. 7. The growing noise region for 4 kmc does not start until at least four scallop wavelengths past the earliest observed 10.7 kmc noise growth. Moreover, the 4-kmc noise pattern resembles that for 10.7 kmc in many details. (The resemblance in details of noise patterns at nearby frequencies has been remarked before, in connection with Fig. 4.)

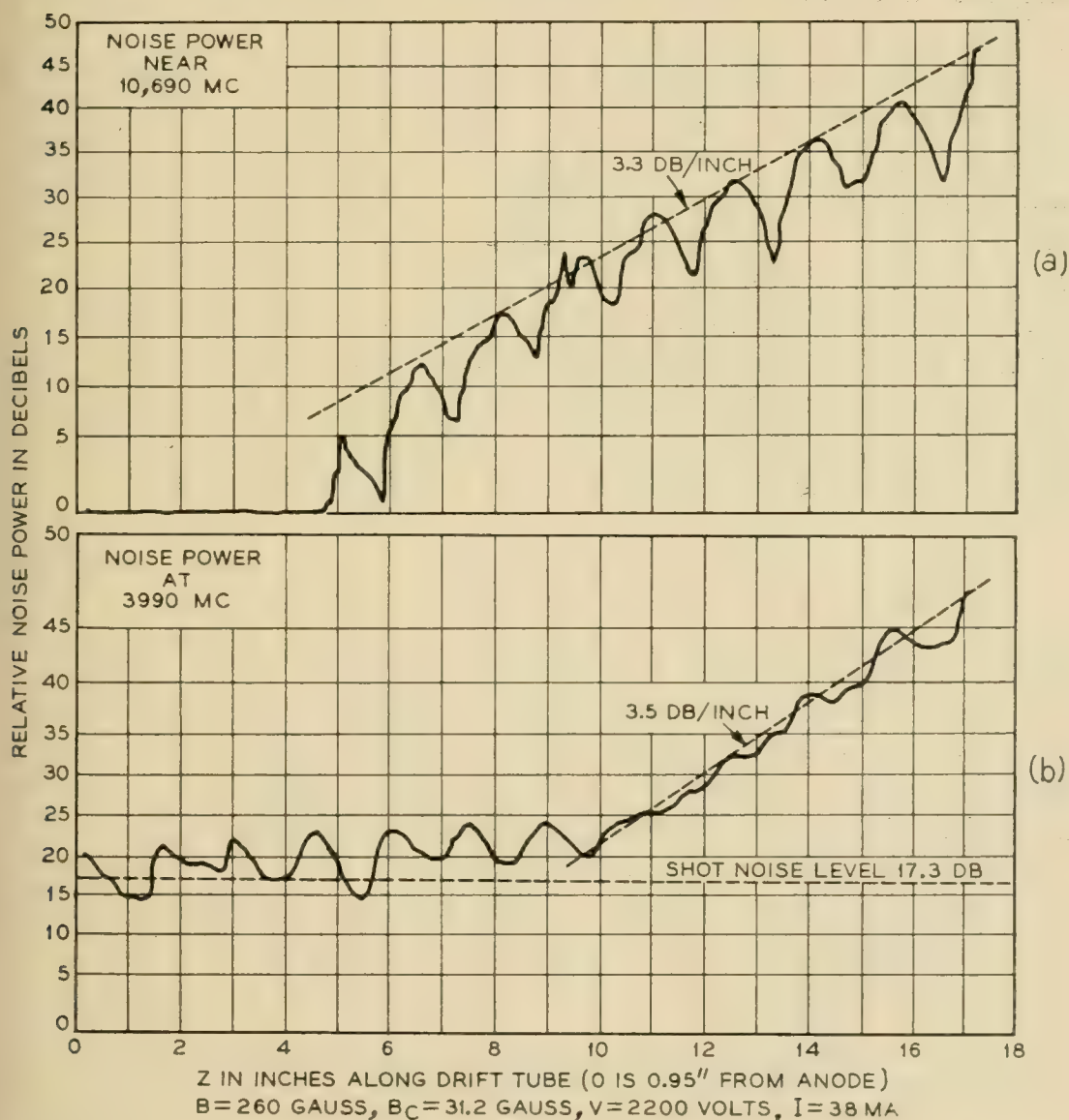


FIG. 7 — The pattern of growing noise in Fig. 6, near 10,690 mcs, is compared with that near 3990 mc for the same fields. The gain region of the latter curve starts much later and is much smoother than the former. The small irregularities on the 3,990 mc curve resemble the scallops of the 10,690 mc curve, in a blurred way.

This information suggests that the growing-noise mechanism is really a two-stage process: amplification of a broad band of microwave frequencies, located far above the observation frequency, followed by a transfer of noise energy to lower frequencies.

(a) *First stage*

At low magnetic fields, there are few ripples per unit length, but their amplitude is usually large; whereas at moderate to large magnetic fields, the ripple amplitude is small, but so is the ripple wavelength. In either case, the bandwidth of RBA is large, usually many thousands of megacycles.

The increase of both bandwidth and gain per scallop with ripple amplitude has been explained by Pierce¹³ by analogy between the gain band of a rippled beam and the stop band of a transmission line filter: a sharply varying periodic disturbance on the latter will reflect short as well as long waves, whereas smoother perturbations will not reflect the shorter waves to any extent.

Another way to study the amplification bandwidth is to derive the equations for RF current in a one-dimensional beam with sinusoidal variation of the reduced plasma wave number, $\beta_q = p \cdot \beta_p$, as Heffner,¹⁴ Bloom,¹⁵ and others^{16, 17} have done. This leads to a Mathieu equation, whose solutions may be studied on the Mathieu stability plot (A, q):

$$\frac{d^2 I}{dx^2} + (A - 2q \cos 2x)I = 0. \quad (6)$$

Here I is the RF current, q a measure of the perturbation amplitude, $x = \pi z/L$, and $A = (2L/\lambda_q)^2$, where L is the ripple wavelength and λ_q the reduced plasma wavelength. Bloom has shown that, if n is the integral number of scallop wavelengths between initial and final planes, the greater the product nq , the greater the total amplification or deamplification, and the less critical the phase relation between RF standing wave and ripple for amplification. Ultimately, for very large nq , amplification will take place for all values of this phase angle.

At higher magnetic fields, both the ripple amplitude and ripple wavelength are decreased. This means that q is reduced, but n increased over any fixed span. This combination usually tends to increase the product nq up to some fairly high field, after which it may decline. More important, the reduction in ripple wavelength shifts the band of amplification to higher frequencies, and greatly increases the frequency band. This occurs because the "resonant" space-charge wavelength is shorter,

and short space-charge wavelengths correspond to high frequencies, where the former change very slowly with frequency.

(b) *Second stage*

When noise power over this large band has been amplified sufficiently, electron bunching becomes non-sinusoidal, and the beam becomes non-linear. Harmonics and beat-frequencies¹⁸ of the fundamentals, and possibly sub-harmonics,¹⁹ are excited. As the beat-frequencies are excited by a continuum of pairs of frequencies, their standing-wave patterns overlap one another, resulting in a "wash-out" of the noise minima, and a smooth growing-noise pattern. Eventually, the same non-linear processes take place within this subsidiary band, leading to a gradual leveling of the entire noise spectrum. The *initial* rates of rise of the intermodulation products, however, should take place closer to the gun and be greater, for a *lower* frequency. They will depend on both the spectrum of noise power in the primary band and, so to speak, the spectrum of "beam non-linearity" within that band.

To simulate this intermodulation process, two low-level klystron signals (9,050 and 12,275 mc, respectively) were simultaneously permitted to modulate the electron beam as it entered the drift tube, by means of a short length of lossy helix. The magnetic fields at the cathode and in the drift space were adjusted to produce a beam ripple which amplified both of these signals simultaneously over most of the drift space, as shown in Fig. 8 (a, b, c).^{*} Noise-power records were then made at the difference-frequency, in the presence of the two modulation signals, Fig. 8(d), and in their absence, Fig. 8(e). The difference between the noise levels in the latter two records increases with distance, as both parent space-charge waves grow in amplitude, and the degree of beam non-linearity increases. Naturally, the contribution to 3,295-mc noise in the absence of modulating signals is far greater than that of the latter two alone, as the primary bandwidth of noise amplified is very great, and that of the signals very small.

There are several reasons why exponentially-growing noise should stop growing and level off, and sometimes even decrease slightly:

- (1) depletion of dc kinetic energy in the beam ripples;
- (2) de-amplification in the fundamental band, due to departure from the proper phase relation for gain between standing wave and ripple, if only over part of the band (Fig. 3);

^{*} The fine ac detail superimposed on the pattern of Fig. 8(b) is due to interference between the waves traveling along the beam and that propagating as a waveguide mode in the drift space.

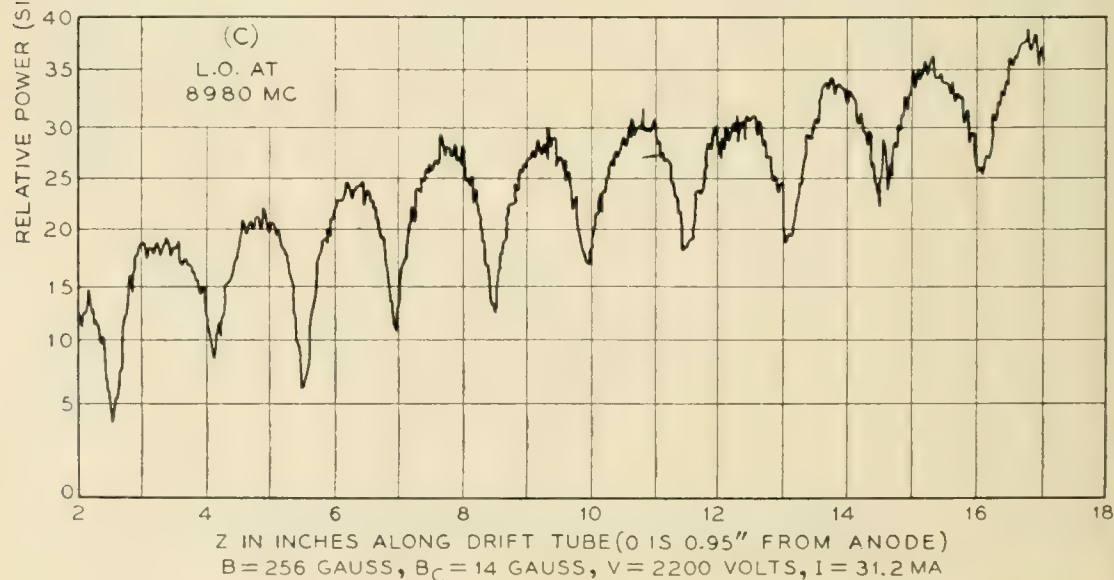
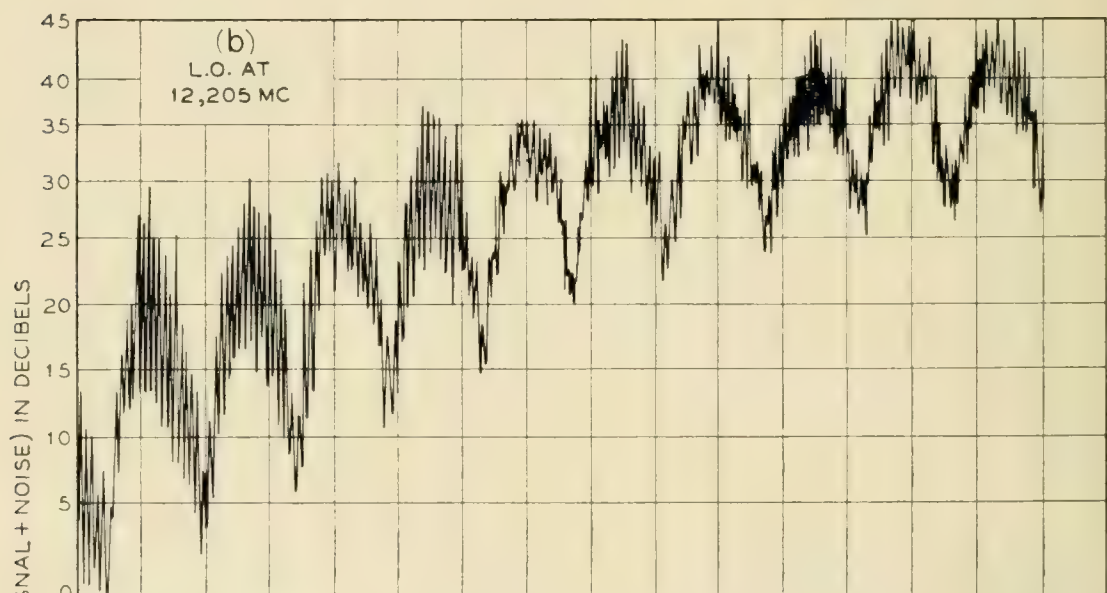
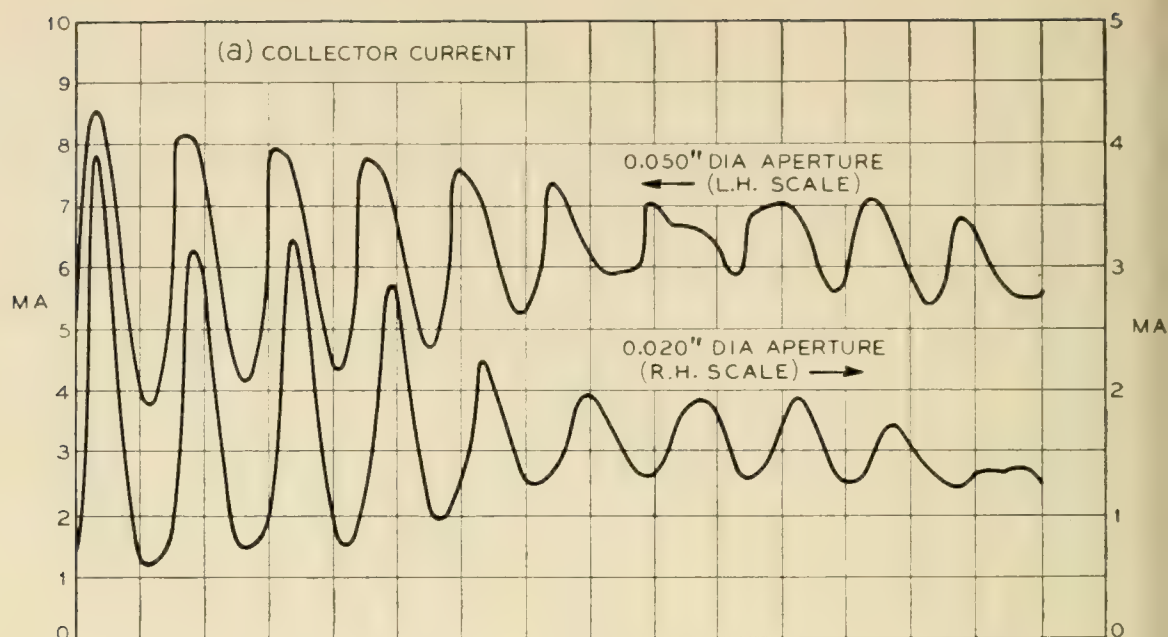


FIG. 8 (a), (b), (c) — The fields have been adjusted to give rippled-beam amplification of two weak klystron signals, 12,275 and 9,050 mc, simultaneously impressed on the beam at the entrance to the drift region.

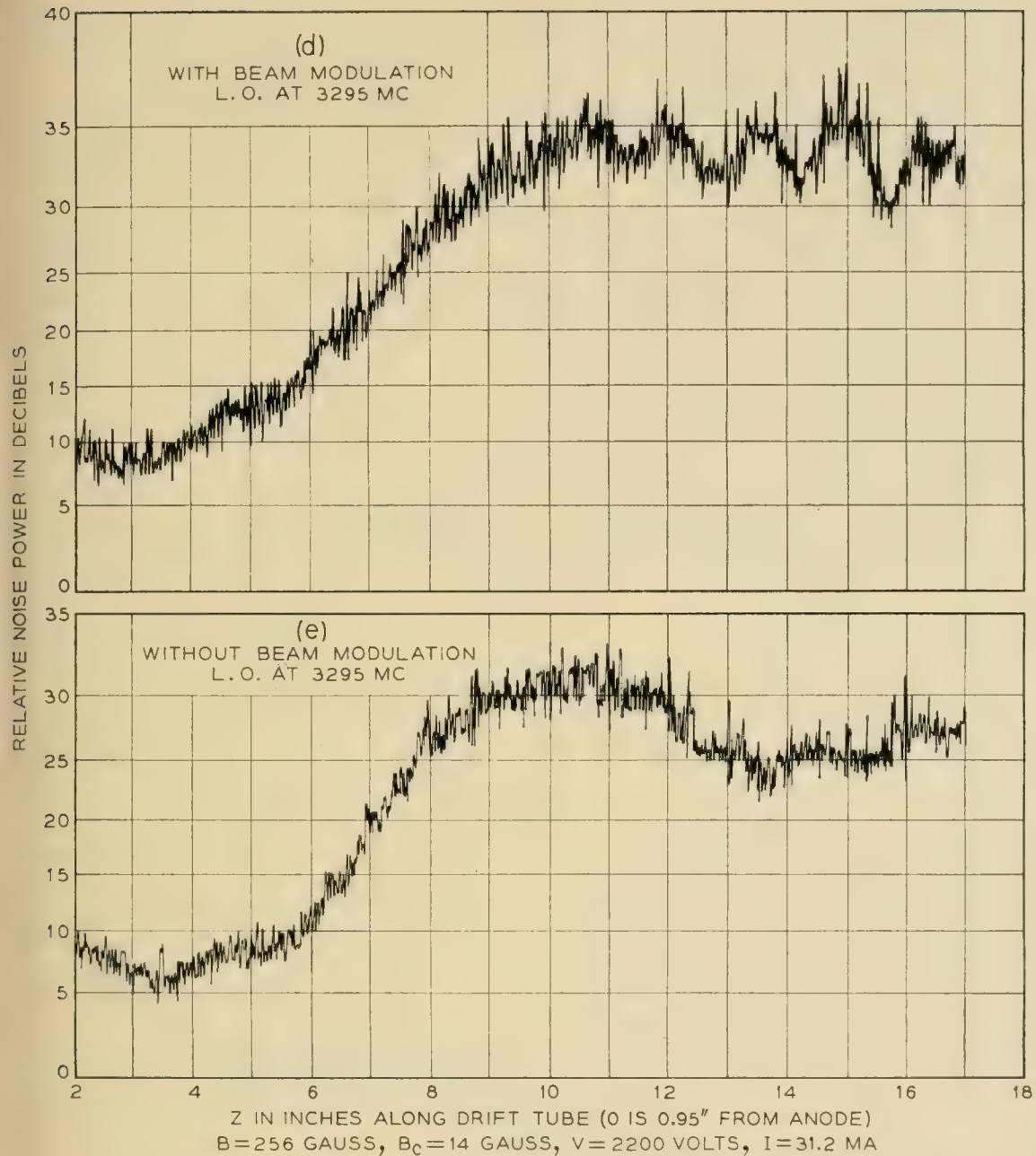


FIG. 8 (d), (e) — Noise power at the difference-frequency, 3,225 mc, is recorded, with and without the signals (b) and (c) present. The difference in ordinates of the two curves increases with distance from the gun, as the impressed signals grow and the beam becomes more non-linear.

(3) sufficient phase shift between inner and outer ripples in the beam for one to de-amplify as much as the other amplifies; and

(4) interference among the intermodulation products excited at different positions along the beam.*

The last effect is illustrated in Fig. 9, in exaggerated form. The beam

* Suggested by C. C. Cutler of Bell Telephone Laboratories.

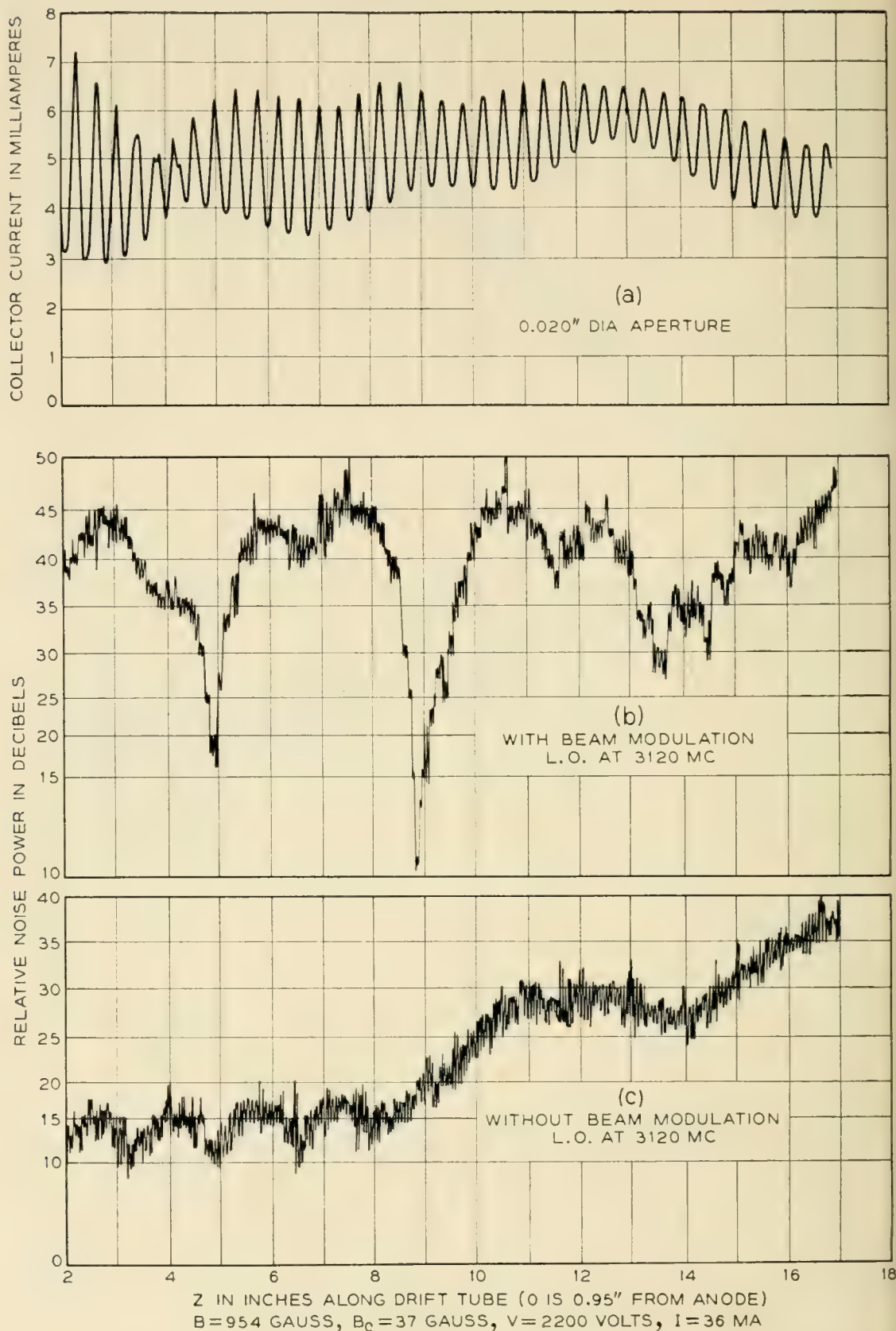


FIG. 9 — Two strong klystron signals are impressed on the beam as in Fig. 8, and noise power at their difference-frequency recorded, both with and without these signals present. The deep noise minima in (b) are due to destructive interference between trains of waves excited by intermodulation at different positions along the beam.

is simultaneously modulated as before with two klystron signals (8,400 and 11,590 mc, respectively), but now at fairly high level; and the focusing field is made large. The interference dips in the pattern of 3,120-mc noise are quite deep, and are spaced irregularly and farther apart than the space-charge wavelength of any of the three frequencies involved. The third dip is shallower than the previous two because of the growth of 3,120-mc noise other than that due to the signals, as shown in Fig. 9(c). The latter pattern of noise in the absence of the two high-level, high-frequency signals suggests that the characteristic first gentle dip following the growing-noise region is indeed of the same nature as the artificially-produced interference dips, and has nearly the same quasi-period.

The pattern of dips agrees with simple calculations, based on this model, in which the amplitude of the difference-frequency intermodulation product, excited at any plane ζ , is assumed proportional to the product of the amplitudes of the two high-frequency space-charge standing waves, as follows:

$$|di_3| \propto |i_1(\zeta) \cdot i_2(\zeta) d\zeta|, \quad (7)$$

where

$$i_n = I_n \sin p_n \beta_p \zeta \cdot \sin \omega_n(t - \zeta/u), \quad (n = 1, 2).$$

The total current at $\zeta = z$ is the sum of contributions from all the standing waves excited to the left of it:

$$|i_3| \propto \frac{1}{4} I_1 I_2 \int_0^z \cos p_3 \beta_p (z - \zeta) [\cos (p_1 - p_2) \beta_p \zeta - \cos (p_1 + p_2) \beta_p \zeta] d\zeta. \quad (8)$$

This expression is readily integrated and evaluated.

VI CONCLUSIONS

Synchronized measurements of electron-current density and noise currents at several microwave frequencies have shown that the "growing noise" pattern in drifting cylindrical beams is the result of a two-stage process. In the first stage, rippled-beam amplification of noise fluctuations takes place over a very broad band of microwave frequencies, much higher than the usual observation frequency. In the second, noise energy is transferred to lower frequencies by intermodulation and other non-linear processes within this band. The element of non-linearity is supplied when primary noise gain is sufficient to make electron bunching

non-sinusoidal. Other sources of non-linearity are thermal velocities, non-laminar beam flow, etc. As the beat-frequency noise increments at any plane are produced by continuous arrays of frequency pairs, increasing in numbers and amplitude in various ways as primary amplification proceeds, the multiple standing-wave patterns at the observation frequency progressively overlap one another. This results in the smooth steep rise of noise power usually observed.

Phase correlation among the space-charge waves excited at successive planes on the beam by the same set of frequency pairs is indicated by gentle dips, due to their destructive interference, in the plateau following the initial noise rise.

Rippled-beam amplification occurs whenever the ripple wavelength and half the space-charge wavelength are nearly equal, and bear a favorable spatial relation to each other. However, this "phase" relation becomes less critical with an increase in either the number of ripple wavelengths over which synchronism persists, or the ripple amplitude, or both. Noise amplification by this mechanism, therefore, is probably present to some degree in all rippled streams, particularly at high fields. The extreme difficulty encountered in focusing ripple-free beams from convergent, shielded guns has to this date prevented the detection of any other primary gain mechanism, which may conceivably co-exist in such beams.

A conspicuous feature of rippled-beam amplification is the decrease in ripple amplitude due to conversion of dc into ac kinetic energy. Such changes in beam structure emphasize the inadequacy of beam-flow computations based entirely on dc force equations. A more detailed description of this dc-ac energy conversion is given in Part II.

ACKNOWLEDGMENTS

The experimental apparatus could not have been built without the combined efforts of many associates of the writer, principally A. R. Strnad, P. Hannes, J. S. Hasiak and J. M. Dziedzic. The author is also indebted to R. Kompfner, C. F. Hempstead and K. M. Poole for valuable suggestions; and above all to C. F. Quate for constant encouragement and advice.

REFERENCES

1. C. C. Cutler and C. F. Quate, Experimental Verification of Space-Charge and Transit Time Reduction of Noise in Electron Beams, *Phys. Rev.*, **80**, p. 875, 1950.
2. L. D. Smullin and C. Fried, Microwave Noise Measurements on Electron Beams, *Trans. I.R.E.* **ED-1**, No. 4, p. 168, Dec., 1954.

3. C. Fried, Noise in Electron Beams, Tech. Rep. 294, Research Laboratory of Electronics, M.I.T., May 2, 1955.
4. J. R. Pierce and L. R. Walker, Growing Waves Due to Transverse Velocities, B.S.T.J., **35**, p. 109, Jan., 1956.
5. G. G. Macfarlane and H. G. Hay, Wave Propagation in a Slipping Stream of Electrons: Small Amplitude Theory, Proc. Royal Soc. (B) **63**, p. 409, 1950.
6. C. K. Birdsall, Rippled Wall and Rippled Stream Amplifiers, Proc. I.R.E., **42**, p. 1628, Nov., 1954.
7. R. W. Peter, S. Bloom, and J. A. Ruetz, Space-Charge-Wave Amplification along an Electron Beam by Periodic Change of the Beam Impedance, RCA Rev., **15**, p. 113, March, 1954.
8. T. G. Mihran, Scalloped Beam Amplification, Trans. I.R.E., **ED-3**, No. 1, p. 32, Jan., 1956.
9. R. H. Dicke, The Measurement of Thermal Radiation at Microwave Frequencies, Rev. Sci. Instr. **17**, p. 268, July, 1946.
10. W. W. Rigrod and J. A. Lewis, Wave Propagation Along a Magnetically-Focused Cylindrical Electron Beam, B.S.T.J., **33**, p. 399, March, 1954.
11. G. M. Branch and T. G. Mihran, Plasma Frequency Reduction Factors in Electron Beams, I.R.E. Trans., **ED-2**, No. 2, 3, April, 1955.
12. Informal communication from H. L. McDowell.
13. Informal communication from J. R. Pierce.
14. Informal communication from H. Heffner.
15. S. Bloom, Space-Charge Waves in a Drifting, Scalloped Beam, unpublished RCA Research Laboratories report.
16. O. E. H. Rydbeck and B. Agdur, Propagation of Space-Charge Waves in Guides and Tubes with Periodic Structure, L'Onde Electrique, **34**, p. 499, June, 1954.
17. P. V. Bliokh and Y. B. Feinberg, Space-Charge Waves in Electron Beams with Variable Velocity, Zhurnal Tekhn. Fiziki, **26**, p. 530, March, 1956.
18. C. C. Cutler, The Nature of Power Saturation in Traveling Wave Tubes, B.S.T.J., **35**, p. 841, July, 1956.
19. S. Lundquist, Subharmonic Oscillations in a Nonlinear System with Positive Damping, Quarterly of Appl. Math., **13**, No. 3, p. 305, Oct., 1955.

Noise Spectrum of Electron Beam in Longitudinal Magnetic Field

Part II — The UHF Noise Spectrum

By W. W. Rigrod

(Manuscript received January 21, 1957)

Sharp peaks are found in the UHF spectrum (10 to 500 mc) of an electron beam, emanating from a shielded diode. In the presence of a longitudinal magnetic field, the strongly rippled beam displays an additional set of peaks whose frequencies are proportional to the field strength. The largest of these, just above the cyclotron frequency, is connected with the overlap of a dense cluster of particle orbits, passing close to the beam axis. It can attain amplitudes of 65 db above background noise.

The transverse distribution of UHF noise power is found to agree with that for ideal Brillouin flow, even in rippled beams. With long ripple wavelengths, two noise maxima are found to flank each beam waist. A small-signal wave analysis explains this pattern, and affords some insight into the energy-exchange processes in rippled-beam amplification. The reduction in "growing noise" due to positive ions is attributed to increased cancellation of net radial beam motion, due to overlap in particle orbits near the axis.

I INTRODUCTION

The reader is referred to Part I¹ for a description of the experimental apparatus and its operation. In this paper, measurements of noise power in the same electron beam are described, with frequencies chiefly in the 10- to 500-mc range, and relatively weak magnetic fields. For the UHF measurements, a calibrated coaxial step attenuator and a super-regenerative receiver (the Hewlett-Packard 417-A VHF Detector) are used. Relative noise-power amplitudes at fixed frequencies are measured as before, in terms of changes in attenuation between probe and receiver required to restore constant receiver output. To obtain qualitative information, however, such as the location of noise maxima along the beam, the series attenuation is fixed. The receiver output is amplified, rectified, and per-

mitted to register itself directly on the chart recorder, whose motion is synchronized with that of the probe. Very roughly, the detector output varies as the log of input power.

Measurements are described (a) of the UHF noise spectrum in the beam, just outside the gun anode; (b) of this spectrum at the end of the drift region, in a longitudinal magnetic field; (c) of the noise-power distribution along the axis; and (d) transverse to the axis of the rippled beam in the drift region. Two calculations are then outlined, one of wave propagation along the rippled beam (to explain the observed distribution patterns), and the other to account for some spectacular peaks in the beam spectrum (b).

II FIELD-INDEPENDENT PEAKS

When the noise spectrum of an electron beam is scanned by a tunable receiver, it is found that an irregular array of narrow-band peaks characterize the UHF region, below about 1000 mc. Of these peaks, some are due to spurious modulation effects,² and can be eliminated as follows:

(1) Transit-time oscillations due to positive ions, secondary electrons, or both. Such frequencies vary with probe (collector) position.

(2) Resonances in the probe and receiver, excited by the pulsed-voltage supply. These are unaffected by changes in collector current.

(3) Ion oscillations in the electron gun or beam. Their frequencies vary with anode voltage.

The remaining narrow-band peaks fall into two classes, depending on whether their frequencies vary with the magnetic field.

Well-defined peaks can be detected with the RF probe stationed one inch from the gun anode, with or without any focusing field. When the beam is focused by a longitudinal magnetic field, these disturbances propagate along the beam, and tend to increase in amplitude with distance, but not to change in frequency. A typical set of such frequencies, within the range of the tunable receiver is as follows: 15.9, 24.3, 31.2, 34.0, 48.5, 63.4, 77.0, 108, 151, 166, 270.5, 372 and 481 mc. (During this measurement, the anode voltage was 2,200, and the peak current about 40 ma.)

No consistent relation could be found between these frequencies and either the anode voltage or the cathode temperature, although unmistakable frequency changes did occur when these parameters were manipulated. Failure to establish such a relation may have been due to uncontrolled drift in cathode activity. In any case, the measurements did serve to narrow the field of possible mechanisms, by eliminating the following:

(1) Transverse positive-ion oscillations,³ for which the frequencies vary as the square root of anode voltage.

(2) Transverse electron plasma oscillations (near or beyond the anode), for which the frequencies would be too high.

(3) Longitudinal electron plasma oscillations at the potential minimum, for the same reason (should be near 2,500 mc).

(4) Longitudinal diode oscillations.⁴ When the electron transit angle through the diode is approximately $(n + \frac{1}{4})$ periods, where n is an integer, the real part of the diode conductance becomes negative, permitting oscillations to occur. Again the frequencies of such oscillations would be too high, (2,200 mc and higher) for the gun used, to conform to the observed values.

There is, however, one published theory for which an order-of-magnitude correspondence does exist between the measured and calculated frequencies. Klemperer^{5, 6} has shown that a strip beam tends to break up into clusters of "pencils" at the cathode. He ascribes these to standing waves resulting from transverse oscillations in the space-charge cloud, and offers an expression for the wave velocity in this medium. Application of his formula to the cathode used in the present experiments results in a least frequency of 31.3 mc. Other observers, such as Smyth⁷ and Veith,⁸ have also reported evidence of interaction between electrons in a retarding-field region and RF fields, which may underlie these oscillations.

III FIELD-DEPENDENT PEAKS

With the RF probe stationed ten or more inches from the gun anode, narrow-band peaks can be found in the noise spectrum of the beam. The amplitudes of these peaks increase and their frequencies decrease with decreases in the magnetic field. For each probe position, the process of finding the peak of greatest amplitude involves repeated adjustments of the focusing field, the magnetic field at the cathode, and the receiver frequency.

When the fields have been so optimized, it is found that the probe is located at or near the first beam-diameter minimum, following that at the entrance to the drift space. When the field is doubled, and the "tuning" process repeated, the greatest peak is found to have about twice the frequency of the first, and the probe is found to be located at or near the second beam waist. It is convenient, therefore, to think of these peaks as "proper" frequencies of the $N = 1$, etc., modes of the rippled beam, where N is the number of ripple wavelengths between gun and probe.

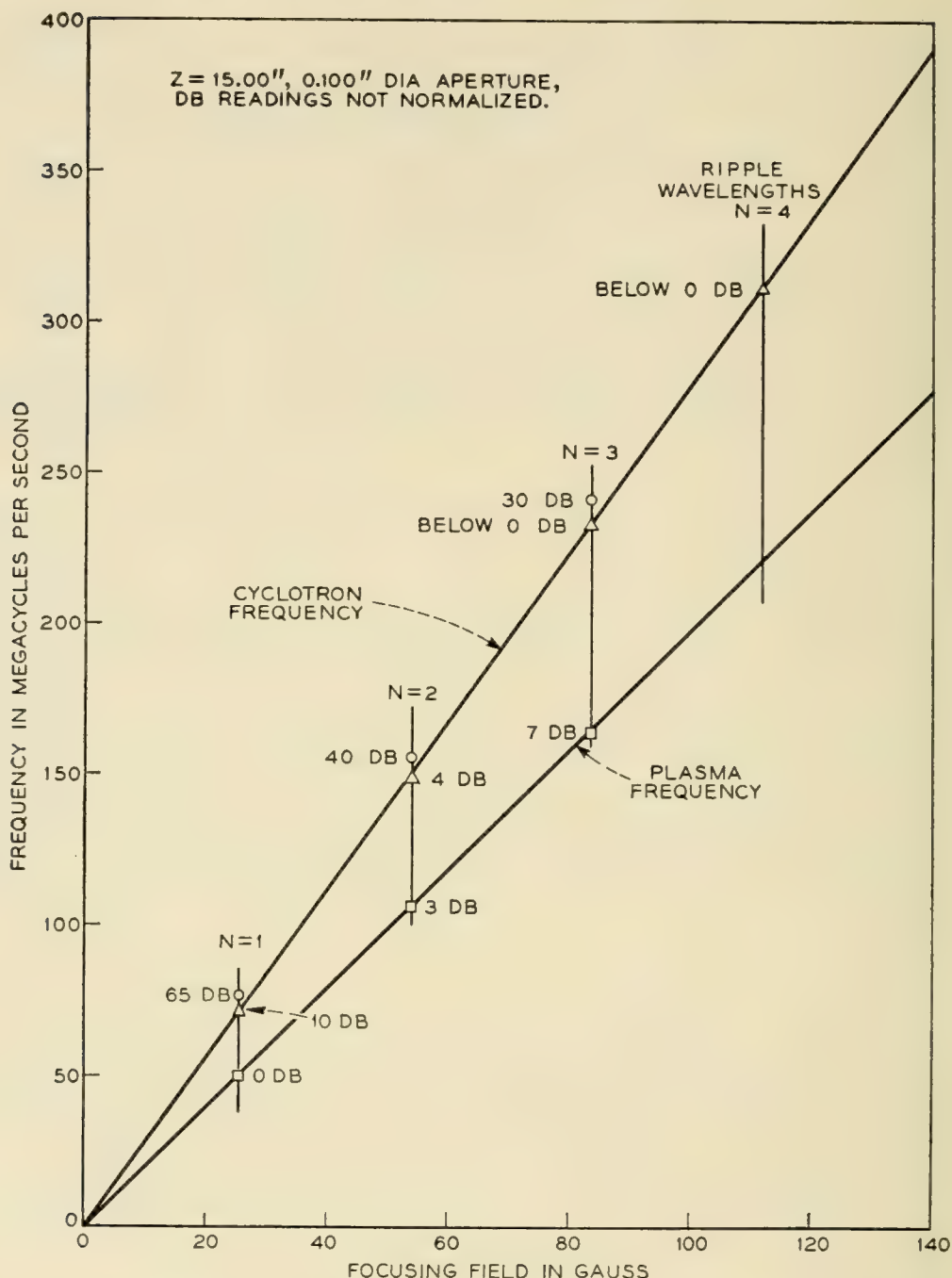


FIG. 1 — Frequencies and amplitudes of several narrow-band UHF peaks measured at a fixed probe position, about 16 inches from the gun anode. N is the number of beam-ripple wave-lengths between anode and probe. Other peaks have been observed at higher harmonics of the "proper" frequency (encircled points), and at about half that frequency.

As shown by the encircled points in Fig. 1, these frequencies range between 1.03 and 1.06 times the calculated cyclotron frequency, and have amplitudes as high as 65 db above the background noise. The amplitudes decrease with increasing N , falling off as the minus two-thirds power of the frequency.

At each of these optimum field settings, several weaker "satellite" peaks can also be detected, most readily those at the cyclotron frequency itself, and at 0.707 times the latter; i.e., the "plasma" frequency, as shown in Fig. 1. In addition, smaller peaks have been repeatedly observed at harmonics (up to the sixth) of the proper frequency, and one at slightly less than half of that frequency. (When a proper frequency was simulated by means of a signal generator, only its first harmonic could be detected in the receiver output.)

At the fields corresponding to $N = 4$ in Fig. 1, the cyclotron frequency (312 mc) was found, but not the proper frequency. The highest proper frequency observed was 240.5 mc, in the $N = 3$ mode. The proper-frequency peaks decrease with increasing focusing field, whereas the field-independent peaks excited in the electron gun tend to increase, at the far end of the drift region.

IV SPATIAL DISTRIBUTION OF UHF NOISE CURRENTS

In Figs. 2 to 5 are shown synchronized chart records of collector current, one or more UHF narrow-band peaks, and microwave noise power near 4,000 mc — all as functions of distance from the electron gun, for the $N = 1$ to 4 modes, respectively. In all runs, the beam was pulsed with a 1,000-cycle square wave, and the collector aperture set at a 0.100 inch diameter. The magnetic fields at the cathode and in the drift space were adjusted before each set of readings, with the probe at a common reference position, for greatest amplitude of some UHF peak. In Figs. 2 and 3, these were proper frequencies, whereas in Figs. 4 and 5 they were field-independent frequencies.

The content of these distribution curves can be summarized as follows:

(1) At the low fields employed (none quite equal to the nominal Brillouin value), the beam ripples are quite large, both in amplitude and wavelength.

(2) The proper-frequency traces have two or three maxima near each beam waist, and their amplitudes grow more rapidly with distance from the gun than any of the satellite frequencies.

(3) The patterns of the cyclotron and "plasma" frequencies do not differ significantly from those of the field-independent frequencies, and usually display two peaks near each beam waist.

(4) The collector-current maxima decrease with distance from the gun, although their minima change little. (The first maximum is sometimes flat-topped due to beam interception before it enters the drift space.) The rate of decrease of these maxima, and the rate of increase of proper-frequency amplitude, are greater, the longer the ripple wavelength.

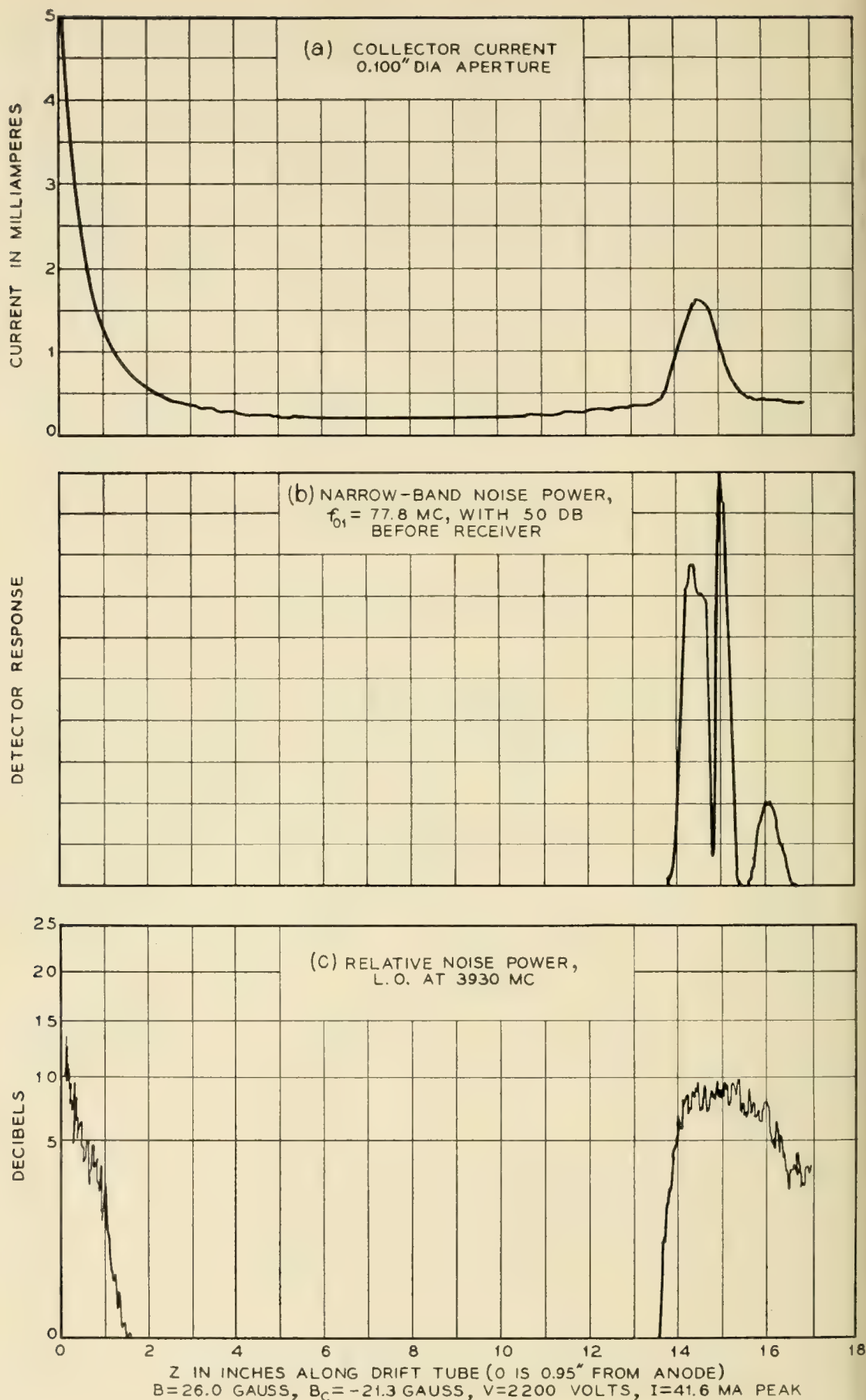


FIG. 2 — The fields have been adjusted for maximum amplitude of the $N = 1$ proper frequency, 77.8 mc, at a reference probe position ($z = 15$ inches). The synchronized probe records indicate three distinct maxima of this proper frequency near the beam waist.

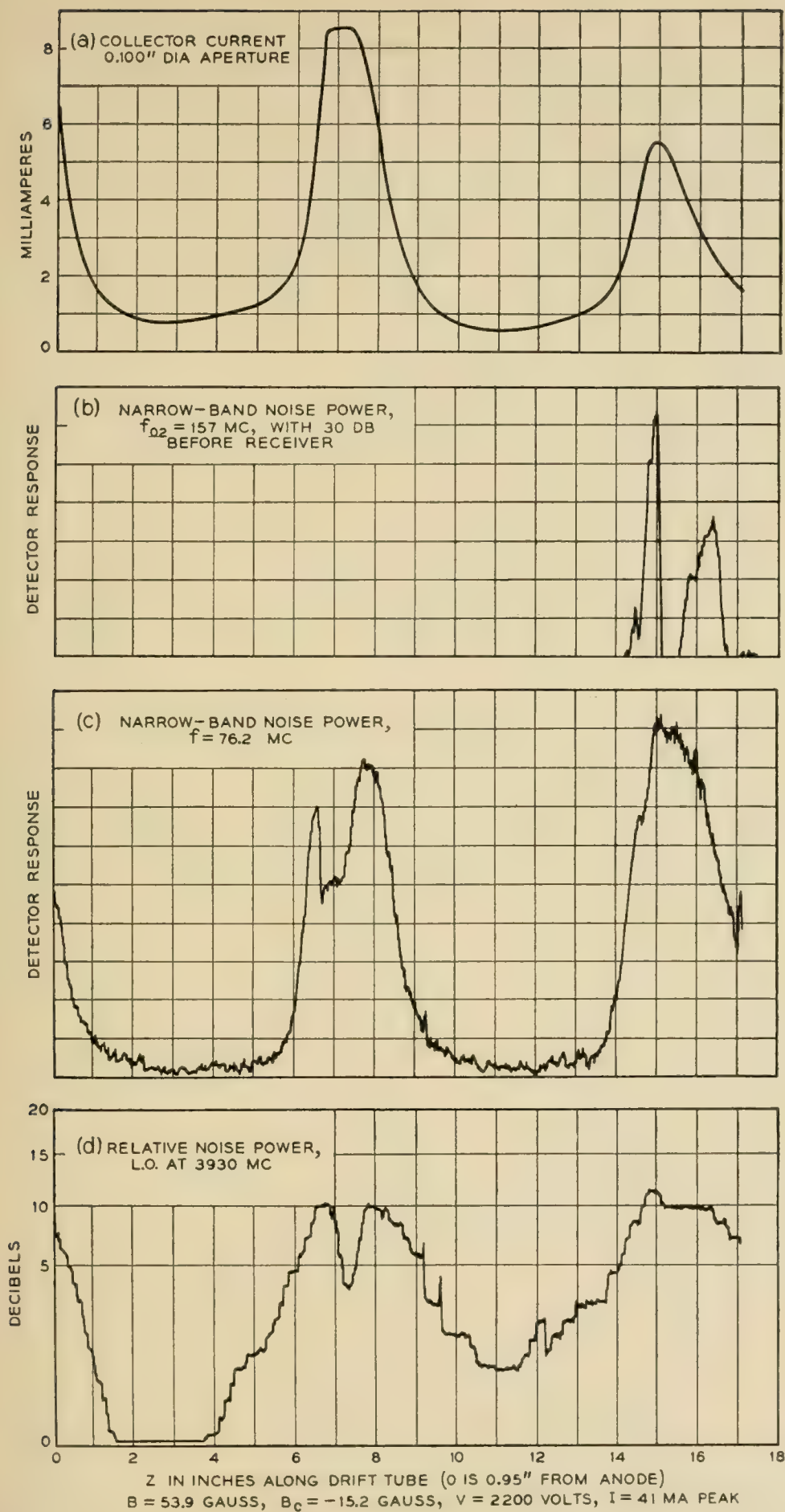


FIG. 3 — The longitudinal distributions of the $N = 2$ proper frequency, 157 mc, as well as its "satellite" 76.2 mc, and microwave noise power, are shown here, with fields adjusted for greatest amplitude of the proper frequency at $z = 15$ inches.

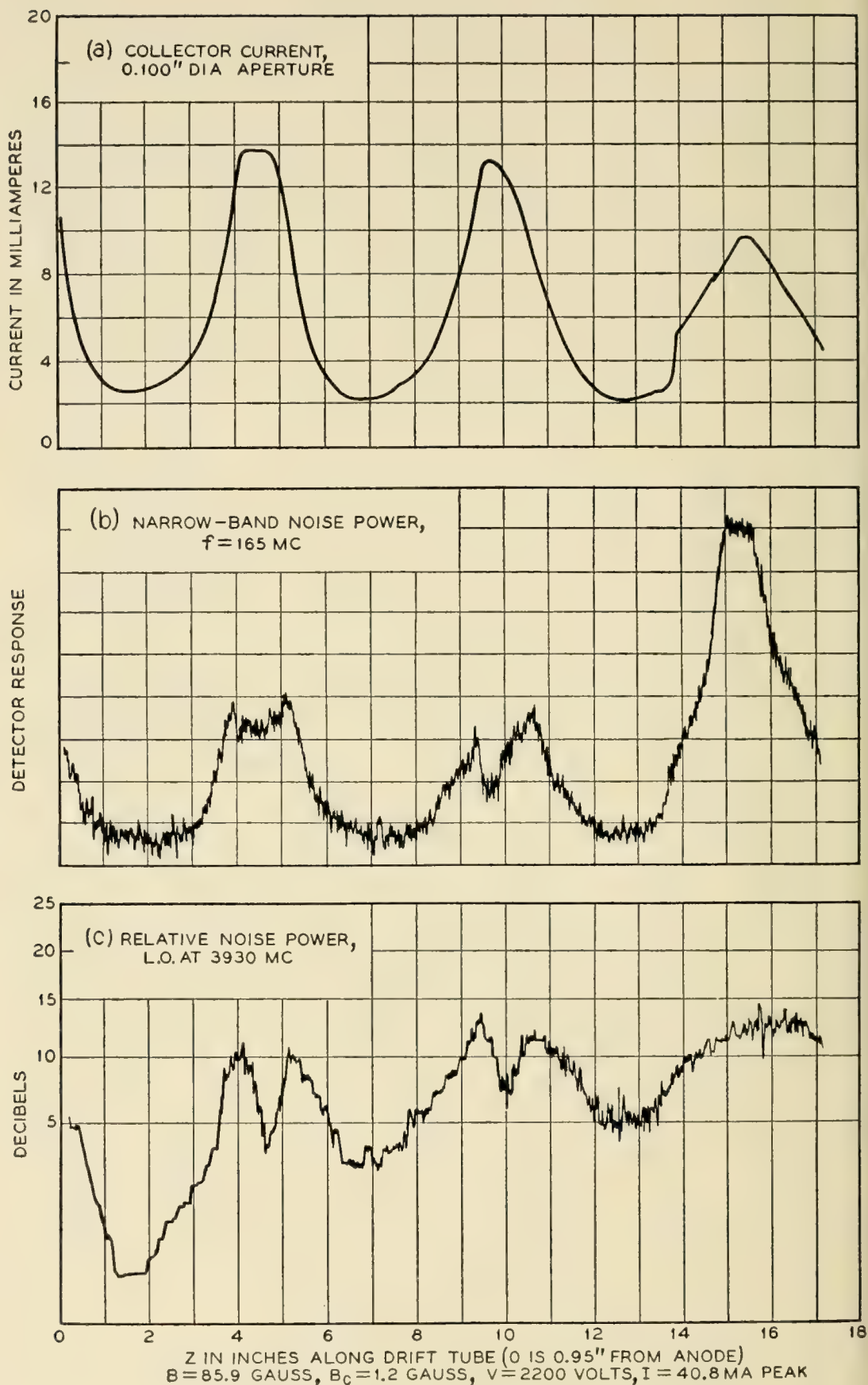


FIG. 4 — The fields have been adjusted for maximum amplitude, at the same reference probe position, of a wave excited in the diode, with frequency unaffected by the magnetic field, 165 mc. The cyclotron frequency for this field is 240.2 mc.

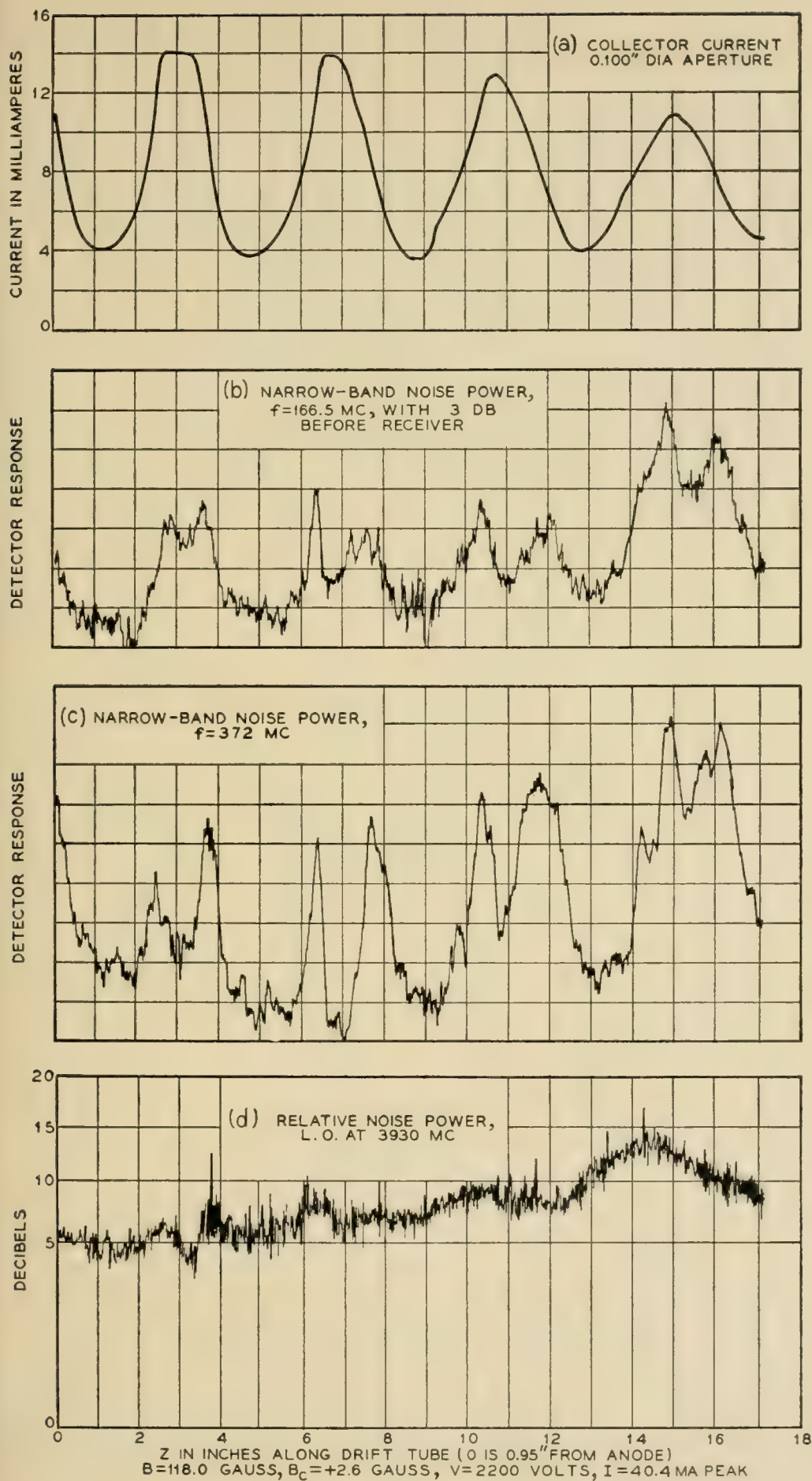


FIG. 5 — A procedure similar to that in Fig. 4 was followed, with four ripple wavelengths between anode and reference plane. The cyclotron frequency here is 329.5 mc.

(5) The patterns of microwave noise power resemble blurred envelopes of the UHF traces.

Some idea of the transverse distributions of UHF noise power and electron-current density, in a region of strong proper-frequency excitation, is given in Figs. 6 and 7. The measurements were taken by moving a small aperture in a broad arc through the probe centerline, just in front of the probe aperture. In both illustrations, the relative noise power has been "normalized" to compensate for variations in electron current traversing the RF gap.

The curves of Fig. 6 are typical of most such measurements. The beam-current density varies smoothly through a single broad maximum, and

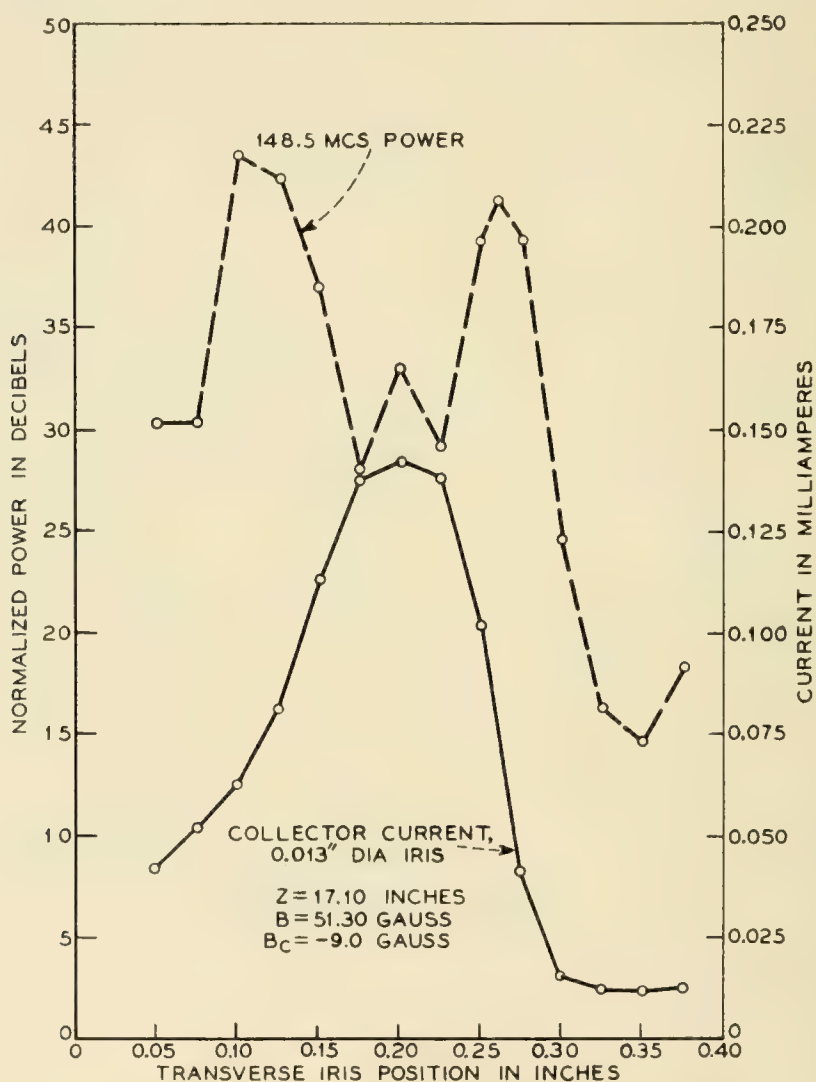


FIG. 6 — Simultaneous point-by-point measurements of collector current and relative noise power, obtained by moving an 0.013-inch diameter aperture in a broad arc through the probe centerline. The probe is stationary, about 18 inches from the gun anode, and the fields have been adjusted for maximum amplitude of the proper frequency, 148.5 mc. The cyclotron frequency is 143.8 mc.

the noise-power density is greatest at the rim of the beam so defined, and least near its center. No evidence of azimuthal periodicity was found. The curves of Fig. 7, which are less typical, indicate five distinct peaks of RF power, despite a nearly symmetrical pattern of collector current. At the time of this measurement, cathode emission may have been uneven, due to coating damage by ion bombardment.

In the rippled beam on which these measurements were made, the ratio of flux encircled at the cathode, to that in the drift space, was very small for most electrons. One would, therefore, expect the transverse noise-power distribution in this beam to resemble that in a smooth Brillouin beam.⁹ The noise power expected when a pinhole aperture is located at the beam center can be compared with that when the aper-

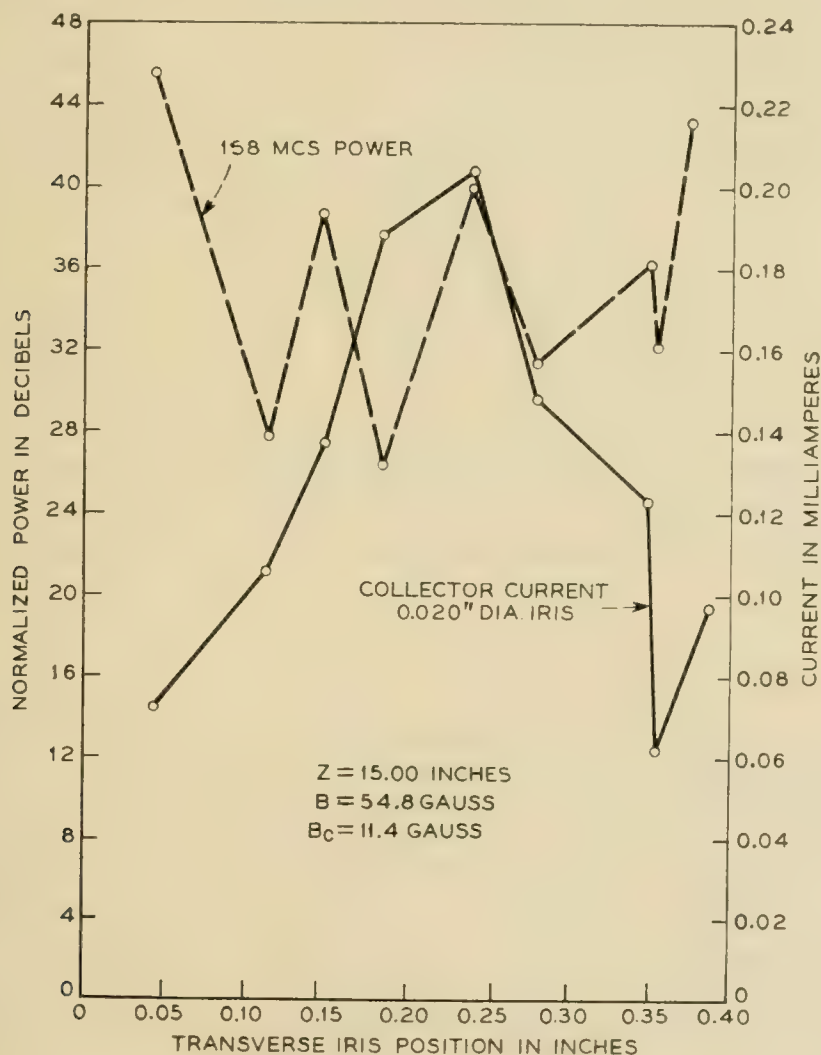


FIG. 7 — Transverse distribution measurements similar to those of Fig. 6. This pattern was obtained a week later than that of Fig. 6, and the cathode was operated at a higher temperature. The cyclotron frequency would be 153.2 mc for the field used.

ture straddles the beam rim, by taking the beam area exposed in the first case to be that of a sector of angle θ , and the length of beam surface in the second case to be that of the corresponding arc:

$$\frac{\text{Rf current sample inside of beam}}{\text{Rf current sample at rim of beam}} \cong \frac{\theta b^2 J_z}{2\theta b G_z} = \frac{b(\omega - \beta u) \cdot I_0(\beta b)}{2u I_1(\beta b)}.$$

Here b is the beam radius, and J_z , G_z the longitudinal components of volume and surface current densities, respectively. I_0 and I_1 are modified Bessel functions, β is the propagation constant, u the beam velocity, and ω the radian frequency. For the frequencies and beam radii employed in these measurements, this ratio is very much less than unity. Thus the pattern of Fig. 6 is in accord with this mode distribution. The multiple peaks of Fig. 7, however, do not conform to this picture, and are not understood at present.

As most of the RF power is concentrated near the rim of the beam, the question arises whether the double and triple peaks, in the longitudinal distribution patterns of Figs. 2 to 5, are not due to the probe aperture breaking through the beam rim. However, the dip between adjacent noise peaks is too great to be explained on the basis of reduced partition noise or weakened gap coupling, assuming the beam diameter there to be less than the gap diameter (0.100 inch). Moreover, double peaks occur even when the beam diameter exceeds the RF gap diameter; for instance, near the last three beam waists of Fig. 5. (When all of the beam is transmitted by the 0.100 inch aperture, the collector-current peak is flat-topped.) It seems likely, therefore, that the double and triple peaks correspond to peaks of amplitude over the entire beam cross-section.

V PROPAGATION ALONG THE RIPPLED BEAM

To find an explanation for the multiple peaks of space-charge current, a small-signal, slow-wave analysis of wave propagation along the rippled beam can be made, in which the special features of these experiments are exploited: long ripple wavelength, effectively no flux at the cathode, and low frequencies. The first of these features suggests that the propagation constants can be evaluated at each cross-section plane as though the beam were uniform, despite the presence of radial velocities. In addition, the space-charge density is assumed constant at each cross-section, and the electron flow laminar.

With these assumptions, the beam can be regarded as a fluid of moving charge, with a single-valued velocity at each point in space, as follows:

$$\underline{v}_0 = (v_r, v_\theta, v_z) \quad (1)$$

where

$$v_r = r \cdot f(z), \quad \text{or} \quad \frac{\partial v_r}{\partial r} = \frac{v_r}{r}, \quad (2)$$

$$v_\theta = r\dot{\theta} = r \frac{\omega_c}{2}, \quad (3)$$

$$v_z = u. \quad (4)$$

Here, r, θ, z are the polar cylindrical coordinates, $\omega_c = \eta B$ the angular cyclotron frequency corresponding to the longitudinal focusing field B , and $f(z)$ a function describing the amplitude and spatial periodicity of the beam ripple. The experimental data indicates that the potential variations along the beam axis are negligible, permitting the assumption that the longitudinal velocity, u , is constant. MKS units are used.

Consistent with the distribution pattern of Fig. 6, the ac field can be represented by an axially-symmetric potential function, similar to that for the smooth Brillouin-flow beam:

$$V \sim I_0(\gamma r) \exp j(\omega t - \beta z), \quad (5)$$

$$\underline{E} = -\text{grad } V. \quad (6)$$

The ac equations of fluid motion are obtained by adding a small ac increment to each of the steady-state velocity components. In addition to the space-charge field, the ac electric field contributes forces acting on the charged medium; those contributed by the ac magnetic field are neglected:

$$\frac{d}{dt} (\underline{v}_0 + \underline{\tilde{v}}) = -\eta [-\text{grad } V - \text{grad } V_0 + (\underline{v}_0 + \underline{\tilde{v}}) \times \underline{B}]. \quad (7)$$

The ac velocity is distinguished by a tilde, and the dc velocity by a zero subscript. Here $\eta = e/m$ is the charge-mass ratio of the electron, a positive quantity. As all ac quantities are functions of spatial positions, their time differentiation (indicated by a dot) is equivalent to multiplication by $j(\omega - \beta u)$, written $j\omega_b$ for brevity.

The components of the force equation are expanded as follows:

$$\begin{aligned} \frac{\partial \tilde{v}_r}{\partial t} + (v_r + \tilde{v}_r) \frac{\partial}{\partial r} (v_r + \tilde{v}_r) + (u + \tilde{v}_z) \frac{\partial}{\partial z} (v_r + \tilde{v}_r) - \frac{(v_\theta + \tilde{v}_\theta)^2}{r} \\ = \eta \frac{\partial V}{\partial r} + \eta \frac{\partial V_0}{\partial r} - \omega_c (v_\theta + \tilde{v}_\theta), \end{aligned} \quad (8a)$$

$$\left[j\omega_b + v_r \left(\frac{\partial}{\partial r} + \frac{1}{r} \right) \right] \tilde{v}_r + \left(\frac{\partial v_r}{\partial z} \right) \tilde{v}_z = \eta \frac{\partial V}{\partial r}, \quad (8b)$$

$$\begin{aligned} \frac{\partial \tilde{v}_\theta}{\partial t} + (v_r + \tilde{v}_r) \left(\dot{\theta} + \frac{\partial \tilde{v}_\theta}{\partial r} \right) + u \frac{\partial \tilde{v}_\theta}{\partial z} + (v_r + \tilde{v}_r) \left(\dot{\theta} + \frac{\tilde{v}_\theta}{r} \right) \\ = \eta \frac{\partial V}{\partial \theta} + \omega_c (v_r + \tilde{v}_r), \end{aligned} \quad (9a)$$

$$\left[j\omega_b + v_r \left(\frac{\partial}{\partial r} + \frac{1}{r} \right) \right] \tilde{v}_\theta = 0, \quad (9b)$$

$$\frac{\partial \tilde{v}_z}{\partial t} + v_r \frac{\partial \tilde{v}_z}{\partial r} + u \frac{\partial \tilde{v}_z}{\partial z} = \eta \frac{\partial V}{\partial z}, \quad (10a)$$

$$\left[j\omega_b + v_r \frac{\partial}{\partial r} \right] \tilde{v}_z = -j\eta\beta V. \quad (10b)$$

An expression for the ac space-charge density, $\tilde{\rho}$, can be obtained in terms of its steady-state counterpart, ρ_0 , by means of the charge-conservation equation:

$$\begin{aligned} \frac{\partial \tilde{\rho}}{\partial t} = -\underline{v}_0 \cdot \text{grad } \tilde{\rho} - \underline{\tilde{v}} \cdot \text{grad } \rho_0 - \rho_0 \text{div } \underline{\tilde{v}} - \tilde{\rho} \text{div } \underline{v}_0 \\ - \underline{v}_0 \cdot \text{grad } \rho_0 - \rho_0 \text{div } \underline{v}_0. \end{aligned} \quad (11)$$

As the beam diameter changes slowly, the dc space-charge density at each plane is taken to be inversely proportional to the square of the radius, b :

$$|\text{grad } \rho_0| = \frac{\partial \rho_0}{\partial z} = -\frac{2\rho_0}{b} \frac{\partial b}{\partial z} \cong -\frac{2v_r}{ur} \rho_0, \quad (12)$$

$$\text{div } \underline{v}_0 = -\frac{\underline{v}_0 \cdot \text{grad } \rho_0}{\rho_0} \cong -\frac{2v_r}{r}, \quad (13)$$

$$\left[j\omega_b + v_r \left(\frac{\partial}{\partial r} + \frac{1}{r} \right) \right] \tilde{\rho} = -\rho_0 \left[\frac{1}{r} \frac{\partial}{\partial r} (r\tilde{v}_r) - j\beta\tilde{v}_z \left(1 - j\frac{2}{\beta u} \frac{v_r}{r} \right) \right]. \quad (14)$$

At low frequencies (the UHF region),

$$\frac{\partial V}{\partial r} \ll \frac{V}{r}$$

as $(\gamma r)^2 \ll 1$. This inequality is also true of other ac quantities proportional to V , such as $\tilde{v}_z$ and $\tilde{\rho}$, and with a small error can be assumed to be true for $\tilde{v}_r$. When the operator $\partial/\partial r$ is omitted from (8), (9), (10), and (14), it is possible to solve explicitly for $\tilde{\rho}$ in terms of ρ_0 and V .

The laminar-flow rippled beam can be described by the particle trajectories, as follows:

$$r_i = r_{0i}(1 + \delta \cos \beta_c z), \quad (15)$$

where r_{0i} is the maximum radius for the particle considered, and $0 < \delta < 1$. For this model of the beam,

$$\frac{v_r}{r} = \frac{\dot{b}}{b} = \frac{-\delta \omega_c \sin \beta_c z}{1 + \delta \cos \beta_c z}, \quad (16)$$

$$\frac{1}{r} \frac{\partial v_r}{\partial z} = \frac{-\omega_c \beta_c \delta (\delta + \cos \beta_c z)}{(1 + \delta \cos \beta_c z)^2}. \quad (17)$$

The region of interest, judging from the observed peak locations, is not at the mid-plane of the beam waist, where $v_r = 0$, but on either side of that plane, where $|v_r/r|$ is greatest. It is readily found that, at these positions $(1/r) (\partial v_r/\partial z)$ is zero, and (14) can be written

$$\tilde{\rho} = \frac{\eta \rho_0 V}{\omega_b^2} \left[\frac{\gamma^2}{\left(1 - j \frac{v_r}{\omega_b r}\right)^2} - \beta^2 \frac{\left(1 - j \frac{2}{\beta u} \frac{v_r}{r}\right)}{\left(1 - j \frac{v_r}{\omega_b r}\right)} \right]. \quad (18)$$

This can be combined with Poisson's Equation,

$$\Delta V = (\gamma^2 - \beta^2)V = -\tilde{\rho}/\epsilon, \quad (19)$$

to furnish a relation between γ and β :

$$\gamma^2 \left[1 - \frac{R}{\left(1 - j \frac{v_r}{\omega_b r}\right)^2} \right] = \beta^2 \left[1 - \frac{R \left(1 - j \frac{2}{\beta u} \frac{v_r}{r}\right)}{\left(1 - j \frac{v_r}{\omega_b r}\right)} \right] \quad (20)$$

where $R = \omega_p^2/\omega_b^2$ and $\omega_p^2 = -\eta \rho_0/\epsilon$, the square of the angular plasma frequency.

At the beam boundary, $r = b$, the continuity of the tangential field components and the change in radial electric displacement can be expressed in the form of an admittance equation:

$$\left[\frac{1}{V} \left(\frac{\partial V}{\partial r} - \frac{\tilde{\sigma}}{\epsilon} \right) \right]_b^I = \left[\frac{1}{V} \frac{\partial V}{\partial r} \right]_b^{II}. \quad (21)$$

Here I refers to the beam, $0 \leq r \leq b$, and II to the space between beam and the concentric conducting tube, $b \leq r \leq a$. The surface charge layer, $\tilde{\sigma}$, takes account of the surface ripple, of amplitude $\tilde{r}$:

$$\begin{aligned}\tilde{\sigma} &= \rho_0 \tilde{r} = -\frac{j\rho_0 \tilde{v}_r}{\omega_b}, \\ -\tilde{\sigma}/\epsilon &= -\frac{R(\partial V/\partial r)}{1 - j\frac{v_r}{\omega_b r}}.\end{aligned}\quad (22)$$

The appropriate potential functions in I and II are reduced by means of the low-frequency, or thin-beam, approximation, as follows:

$$\begin{aligned}\left[\frac{1}{V}\frac{\partial V}{\partial r}\right]_b^{\text{I}} &= \frac{\gamma I_1(\gamma b)}{I_0(\gamma b)} \cong \frac{\gamma^2 b}{2}, \\ \left[\frac{1}{V}\frac{\partial V}{\partial r}\right]_b^{\text{II}} &= \beta \left[\frac{I_1(\beta b)K_0(\beta a) + I_0(\beta a)K_1(\beta b)}{I_0(\beta b)K_0(\beta a) - I_0(\beta a)K_0(\beta b)} \right] \\ &\cong \frac{\beta^2 b}{2} - \frac{1}{b \ln a/b}\end{aligned}\quad (23)$$

where the following small-argument approximations have been used:

$$\begin{aligned}K_0(x) &\cong -\ln x, \\ K_1(x) &\cong \frac{x}{2} \ln x + \frac{1}{x}.\end{aligned}\quad (24)$$

The boundary equation thus provides a second relation between γ and β :

$$\frac{\gamma^2 b}{2} \left[1 - \frac{R}{\left(1 - j\frac{v_r}{\omega_b r}\right)} \right] = \frac{\beta^2 b}{2} - \frac{1}{b \ln \frac{a}{b}}. \quad (25)$$

For the smooth beam in Brillouin flow ($v_r = 0$), the boundary equation, to the same low-frequency approximation, is as follows:

$$R_0 \equiv \frac{\omega_p^2}{(\omega - \beta_0 u)^2} = \frac{2}{(\beta_0 b)^2 \ln \frac{a}{b}}. \quad (26)$$

To see how the beam ripple affects the propagation constant, it is sufficient to find its first-order effect; i.e., to assume relatively small radial velocities and find a solution for β which is not very different from its value, β_0 , for the smooth beam:

$$\beta = \beta_0 + \delta = \beta_e \pm \beta_q + \delta, \quad (27)$$

where

$$|\delta| \ll \beta_0; \quad \beta_q \cong \frac{\omega_p}{u \sqrt{R_0}}; \quad \beta_e = \frac{\omega}{u}.$$

In addition, $|v_r/\omega_b r|$ is less than unity, and $|2v_r/\beta ur|$ can be neglected entirely. With these assumptions, the boundary and characteristic equations can be combined to solve for β :

$$\frac{\gamma^2}{\beta^2} = \frac{F(F - R)}{F^2 - R} = \frac{F - \left(\frac{\beta_0}{\beta}\right)^2 R_0}{F - R}, \quad (28)$$

where

$$F = 1 - j \frac{v_r}{\omega_b r}.$$

Utilizing the low-frequency condition, $|R^2| > |R| > 1$, this equation can be reduced and, after some algebra, solved:

$$\frac{\beta_0}{\beta} \left(\frac{R_0}{R}\right)^{1/2} = \frac{\beta_0(\delta \pm \beta_q)}{(\beta_0 + \delta)(\pm \beta_q)} \cong 1 - j \frac{v_r}{2\omega_b r},$$

$$\beta_s \cong \beta_e + \beta_q - j \frac{v_r}{2ur}, \quad (29a)$$

$$\beta_f \cong \beta_e - \beta_q + j \frac{v_r}{2ur}. \quad (29b)$$

These expressions show that the current in the slow wave (I_s) will grow when (v_r/r) is negative; i.e., when the beam is contracting, and decrease during its expansion. The fast wave (I_f) will do the opposite. In probe measurements along the beam, the detected ac power is proportional to the square of the total space-charge current, which has the following dependence on time and distance when the amplitudes of both waves are initially equal:

$$(I_s + I_f) = 2I_{\max} \cos(\omega t - \beta_e z) \cdot \cos(\beta_q z) \cdot \sinh\left(\frac{V_r z}{2ur}\right). \quad (30)$$

In UHF noise-power measurements along beams with long ripple wavelengths, the two planes of maximum $\pm (v_r/r)$ are separated by only a small fraction of a space-charge wavelength. Therefore, $\cos \beta_q z$ at the first of these planes is only slightly larger than at the second. Thus, two peaks of current are observed, in agreement with (30). By contrast, in rippled-beam amplification at microwave frequencies, shorter ripple wavelengths and smaller ripple amplitudes are employed. Then (v_r/r) varies nearly sinusoidally over the ripple wavelength. For maximum net gain per ripple, maximum negative (v_r/r) is adjusted to coincide with the plane of $\cos \beta_q z = 1$ (maximum current), and maximum positive (v_r/r) at the current minimum, half a wavelength beyond.

The gain constants in (29) are independent of frequency. The net gain per ripple wavelength, however, will vary with frequency, depending on how closely both the current maxima and minima coincide with the regions of maximum $\pm (v_r/r)$, respectively. This is a statement of the "resonance" condition between ripple wavelength and half the space-charge wavelength, which emerges from one-dimensional analyses¹⁰ of this gain mechanism based on transmission-line analogies.

Such analyses generally assume small-amplitude sinusoidal variations of the reduced plasma wave number, β_q , along a one-dimensional beam in a longitudinal ac field with no losses. Periodic variations in either beam or wall diameters, or beam velocity, cause the beam "impedance" to vary periodically, imparting to it narrow-band filter-like properties equivalent to narrow-band signal gain. From another point of view,¹¹ these periodic impedance changes couple the fast and slow space-charge waves to each other intermittently, thereby effecting an energy transfer from the fast to the slow wave. As this coupling is lossless, I_s increases and I_f decreases with drift distance, in such a way as to keep their product constant. Then the product $I_{\max}I_{\min}$ increases, and the ratio $I_{\max}/I_{\min}$ correspondingly decreases. In the case of *noise-power* amplification, two uncorrelated space-charge standing waves are present. Because the two slow waves cannot simultaneously be amplified at the expense of the two fast waves, the product $I_{\max}I_{\min}$ must remain constant.

The observed noise-current patterns in rippled-beam amplification,¹ however, are characterized by a *nearly constant ratio* $I_{\max}/I_{\min}$, and an *increase in the product* $I_{\max}I_{\min}$ along the beam, despite the fact that the beam voltage is fixed. This apparent contradiction can be resolved by a closer look at the energy-exchange processes.

Chu¹² has shown that the kinetic power flow in space-charge waves (the major part of the total power) is equal to the difference in powers carried by the fast and slow waves. This is equally true of beams with transverse motions and fields.¹³ In rippled-beam amplification, whether analyzed as a modulated linear beam or at each beam cross-section separately, as here, the propagation constants are found to be complex conjugate quantities, whose real parts describe the ordinary fast and slow waves of a uniform beam. From either point of view, therefore, a decrease in I_f and an increase in I_s signifies an increase in the negative kinetic power flow carried by the waves, or a decrease in the total kinetic energy of the beam.

As shown in (29), the gain constants are proportional to v_r , indicating that the dc energy transferred to the waves when the beam contracts could only have come from the *radial* kinetic energy, not the longitudinal.

The direction of energy transfer is reversed during the subsequent beam expansion. If the ripple were perfectly symmetrical, therefore, and the dc-ac energy exchange perfectly reversible, the net effect of a beam ripple would be zero. Neither of these conditions is quite true in actual beams. Rippled flow is never truly laminar, and $|v_r/r|$ usually decreases with drift distance as the flow loses coherence; i.e., it is greater in beam contraction than in the next expansion. This by itself would produce a net gain per ripple in I_s , and a net loss in I_f , of equal amounts. In addition, however, unavoidable small non-linearities in electron motions prevent all of the ac energy in a de-amplified wave from being converted back to dc kinetic energy. Thus it is possible for *both* the fast and slow waves to increase in a ripple wavelength, the latter always more than the former.

The greater gain of the slow wave entails a loss of radial kinetic energy, in agreement with the observation that the ripple amplitude always decays more rapidly when rippled-beam amplification takes place. The incomplete reversibility of the ac-dc energy exchange probably accounts for the observed increase in $I_{\max}I_{\min}$ for noise currents. Finally, the net amplification of all of the space-charge waves, fast as well as slow, is in accord with the observed near-constancy of the ratio $I_{\max}/I_{\min}$ for microwave-frequency noise, despite increases in the product $I_{\max}I_{\min}$ of 30 db and more.

VI ORIGIN OF THE PROPER-FREQUENCY PEAKS

Of the various peaks in the beam's noise spectrum, described in Section III and Fig. 1, those with "proper frequencies," slightly above the cyclotron value, are so large in amplitude that even an approximate analysis should be able to account for them. To do so, a "working model" of the beam is needed, which conforms to the experimental conditions which existed during the observations:

- (1) The peak intensities were greatest near the middle of each beam waist, and decreased with decrease in ripple amplitude.
- (2) The focusing field was below the nominal Brillouin value. The field at the cathode, B_c , was finite and opposed to the main field, B .
- (3) Collector-current measurements along the beam axis showed the ratio of maximum to minimum current to be greater, the smaller the aperture.
- (4) The gas pressure was about 10^{-7} mm Hg. The beam was pulsed with a 1,000-cycle square wave.

Item (3) indicates that the flow was non-laminar; and Item (4) indicates the presence of positive ions. All the items are consistent with the following picture:

In a beam with large ripples, nearly all electrons have their maximum radii and zero radial velocity at the same z -plane. Those with sufficiently large maximum radius will have enough transverse kinetic energy to surmount the space-charge forces at the beam waist, and pass through or close to the axis. Others, with smaller maximum radii, will spiral about that axis. Dolder and Klemperer¹⁴ have observed a similar division of electrons into "crossovers" and non-crossovers, in electron-optical systems without magnetic fields.

Positive ions tend to neutralize the electronic space charge at the beam waists, broadening the region in which crossover occurs. The crossover trajectories thereupon overlap one another, resulting in multivalued transverse particle velocities in this region. In a first-order (linearized) study of wave propagation along the beam, one must replace the actual multivelocity charge motions with a single "fluid" of charge, whose velocity at any point is the average of the particle velocities there. It is clear that the z -velocity of the stream is u , and the radial velocity zero. The tangential velocity, $v_\theta = (\dot{\theta}r)_{av}$, however, is more complicated.

Owing to the partial or total neutralization of electronic space charge at the beam waists, and their large radii elsewhere, the crossover electrons will encounter virtually no space-charge forces in their paths. Their transverse paths will consequently be circles about fixed centers, described with angular velocity equal to the cyclotron frequency. Their angular velocity about the beam axis is given by Busch's Theorem:

$$\dot{\theta} = \frac{\omega_c}{2} \left[1 + \frac{K}{r^2} \right],$$

where

$$K = -r_c^2 \left(\frac{B_c}{B} \right) = r_{\max} r_{\min} \quad (31)$$

is a positive quantity, as B_c/B is negative. Here, r_c is the radius at which a particular electron left the cathode, and r is its radius in the drift region. The angular velocity, $\dot{\theta}$, is greater than $\omega_c/2$ at all times, and exceeds ω_c in the waist region of the beam. The average value of v_θ at any point here, therefore, is greater than $\omega_c r$ and presumably varies from point to point in some unknown way.

If v_θ is left unspecified, and the assumptions adopted of zero space-charge forces and radial velocity over a finite length of beam:

$$\eta \frac{\partial V_0}{\partial r} = 0, \quad v_r = 0, \quad \frac{dv_r}{dt} = 0, \quad (32)$$

the radial component of the force equation (7) in Euler coordinates can

be written as follows:

$$\left(\frac{dv_0}{dt}\right)_r = -\frac{v_\theta^2}{r} = -\omega_c v_\theta, \quad (33)$$

$$v_\theta = 0 \quad \text{and} \quad \omega_c r. \quad (34)$$

Thus, the radial "balance" conditions (32) are consistent with either of two values for v_θ , of the equivalent stream with single-valued velocities. As it develops that either of these values leads to the same result, the first one will be used here for simplicity, $v_\theta = 0$.

An ac traveling wave along this beam cannot have any θ -dependency, because the beam has no single value of angular velocity $\dot{\theta}$, which might remain in synchronism with that of the wave. Thus, the perturbed dynamics equation (7) can be expanded, with the assumptions of an axial-symmetric ac field given by (5) and (6), a stream with steady-state velocity $(0, 0, u)$, constant space-charge density ρ_0 , and no space-charge forces, as follows:

$$\frac{\partial \tilde{v}_r}{\partial t} + u \frac{\partial \tilde{v}_r}{\partial z} = \eta \frac{\partial V}{\partial r} - \omega_c \tilde{v}_\theta,$$

$$\frac{\partial \tilde{v}_\theta}{\partial t} + u \frac{\partial \tilde{v}_\theta}{\partial z} = \omega_c \tilde{v}_r,$$

$$\frac{\partial \tilde{v}_z}{\partial t} + u \frac{\partial \tilde{v}_z}{\partial z} = \eta \frac{\partial V}{\partial z}.$$

These are solved for the ac velocity components:

$$\tilde{v}_r = \frac{-j\eta\omega_b}{\omega_b^2 - \omega_c^2} \frac{\partial V}{\partial r}, \quad (35)$$

$$\tilde{v}_\theta = -\frac{j\omega_c}{\omega_b} \tilde{v}_r, \quad (36)$$

$$v_z = -\frac{\eta\beta}{\omega_b} V. \quad (37)$$

With $\text{grad } \rho_0 = \text{div } \underline{v}_0 = 0$, the charge-conservation equation (11) can be solved for $\tilde{\rho}$:

$$\begin{aligned} \frac{\partial \tilde{\rho}}{\partial t} &= -\underline{v}_0 \cdot \text{grad } \tilde{\rho} - \rho_0 \text{div } \underline{\tilde{v}}, \\ \tilde{\rho} &= \frac{j\rho_0}{\omega_b} \text{div } \underline{\tilde{v}} = \eta\rho_0 \left[\frac{\gamma^2}{\omega_b^2 - \omega_c^2} - \frac{\beta^2}{\omega_b^2} \right]. \end{aligned} \quad (38)$$

At very large ripple amplitudes, it is a fair assumption that the density of non-crossover electrons is negligible relative to that of crossovers in this region. Poisson's equation (19) can then be combined with the

above expression to obtain the characteristic equation:

$$\left(\frac{\beta}{\gamma}\right)^2 = \frac{1 - \frac{\omega_p^2}{\omega_b^2 - \omega_c^2}}{1 - \frac{\omega_p^2}{\omega_b^2}}. \quad (39)$$

In a frame of reference moving with the stream, u' is 0, $\beta'u'$ is 0, and $\omega_b' = \omega$. Then,

$$\left(\frac{\beta'}{\gamma'}\right)^2 = \frac{(\omega^2 - \omega_c^2 - \omega_p^2)\omega^2}{(\omega^2 - \omega_p^2)(\omega^2 - \omega_c^2)} \quad (40)$$

and β' becomes an infinite imaginary quantity when ω is ω_c . The phase velocity in the moving frame is infinite, as the real part of β' is zero; therefore the phase velocity v_p in the rest frame is also infinite. Thus, there is no Doppler shift in the "resonant" frequency observed in the rest frame:

$$\omega_{\text{observed}} = \frac{\omega_c}{1 - \frac{u}{v_p}} = \omega_c. \quad (41)$$

As the actual beam has a z -velocity spread, the field is never perfectly uniform, and as the calculation is valid for small ac quantities only, the discrepancy between this result and the observed "proper" frequencies, which were 1.03 to 1.06 times the cyclotron value, is not unexpected.

The singularity in (40) is seen to disappear when $\omega_p^2 = 0$. This indicates that an exact calculation would show the gain constant ($-j\beta$) to increase with ρ_0 , the density of the crossover electrons. Their trajectories, described by (31), and the absence of space-charge forces are such that $K = r_{\text{max}}r_{\text{min}}$; that is, the greater r_{max} , the smaller r_{min} , the distance of closest approach to the axis. Thus, a larger ripple amplitude (permitted by a lower magnetic field) produces a greater electron density in the waist region, and accordingly a greater oscillation amplitude at the resonant frequency, as observed.

The foregoing mathematics describes a form of resonance, the infinite phase velocity corresponding to longitudinal "cutoff" in a waveguide. Unlike a waveguide, however, the disturbance increases rather than attenuates along the axis, due to the transfer of dc kinetic energy (represented by v_0^2/r) to the ac fields (excited by noise fluctuations at the cathode), at the cyclotron frequency ω_c .

Except for the direction of energy transfer, the situation is analogous to that of a low-pressure gas in a uniform magnetic field, when stressed by an impressed ac field of varying frequency. It has been found that the breakdown field at the cyclotron frequency is very much less than

at other frequencies.¹⁵ Here the energy supplied by the ac field is coupled most effectively to the free electrons at the resonant frequency, increasing their dc kinetic energy until the gas breaks down. The circular ac charge motions due to the dc magnetic and the ac electric fields are superimposed on high-velocity random motions, similar to the radial motions in the drifting beam.

The UHF peaks observed at harmonics of the proper frequency may simply be due to the non-linear character of the beam, when excited by the high-level fundamental oscillations. The other faint satellite peaks, near $0.5 \omega_c$ and $0.707 \omega_c$, seem to be associated with the unneutralized space-charge density at the beam waist.

The conspicuous role played by crossover electrons in the waist region of rippled beams, due to the tendency of their orbits to overlap there, leads one to re-examine their influence on rippled-beam amplification. As seen in the previous section, this gain process depends on the average value of (v_r/r) at each cross-section plane of the beam. The fraction of all electrons which penetrate to the beam axis depends on competition between the unneutralized space-charge forces and the particle's transverse kinetic energy. An increase in positive ion density tends to make the potential depressions at beam waists broader and shallower, and thereby increase the number of crossover electrons as well as the axial distance over which they reach the axis. The net effect is to reduce the average value of $|v_r|$ over a greater portion of the ripple wavelength, and thus reduce the net gain of the space-charge wave. This may explain why the "growing noise" phenomenon tends to be inhibited by an increase in positive ion density.

VII CONCLUSIONS

Evidence is found of oscillations with frequencies in the 10- to 500-mc region inside of an electron-gun diode. There is some basis for associating them with electron-field interaction in the retarding region of the diode. Another type of narrow-band noise peak is found near the waists of a strongly rippled beam in a longitudinal magnetic field, with frequencies proportional to the field strength. The strongest of these, at about 1.05 times the angular cyclotron frequency, ω_c , as well as its harmonics, can be explained by the resonant behavior of a short section of the beam, in which the average transverse velocity is nullified by overlap in particle orbits. Fainter satellite peaks, near $0.5 \omega_c$, $0.707 \omega_c$, and ω_c , respectively, accompany the dominant frequency.

In a drifting beam launched from a shielded electron gun and focused by an axial field, the transverse distribution of noise (or signal) intensity is found to agree with that predicted for ideal Brillouin flow. Despite

the presence of thermal motions and beam ripples, the ac power is found to be concentrated chiefly at the rim of the beam. Occasionally, several concentric rings of noise maxima are found within the beam, possibly due to unusual cathode conditions.

When the ripple wavelength is very long, two maxima of noise power are observed to flank each beam waist. A first-order calculation of wave propagation along a rippled laminar-flow beam accounts for this pattern by showing that space-charge waves grow at the expense of dc kinetic energy in the radial charge motion. In rippled-beam amplification of noise, the product $I_{\max}I_{\min}$ has been found to increase, and the ratio $I_{\max}/I_{\min}$ remain nearly constant, because both fast and slow waves are amplified, the former less than the latter, and because the wave coupling is not lossless.

Positive ions tend to collect at the waist of rippled beams, thereby extending the region in which electrons pass close to the axis, instead of circling about it. The overlap of their orbits leads to net cancellation of radial charge motion, and hence a reduction in rippled-beam amplification. This may explain why positive ions tend to inhibit the "growing noise" phenomenon.

REFERENCES

1. W. W. Rigrod, Noise Spectrum of Electron Beam in Longitudinal Magnetic Field. Part I — The Growing Noise Phenomenon, p. 831 of this issue.
2. C. C. Cutler, Spurious Modulation of Electron Beams, *Proc. I.R.E.*, **44**, p. 61, Jan., 1956.
3. K. G. Hernquist, Plasma Ion Oscillations in Electron Beams, *J. Appl. Phys.* **26**, p. 544, May, 1955.
4. F. B. Llewellyn and A. E. Bowen, The Production of UHF Oscillations by Diodes, *B.S.T.J.*, **18**, p. 280, April, 1939.
5. O. Klemperer, Influence of Space Charge on Thermionic Emission Velocities, *Proc. Royal Soc. (London) (A)* **190**, p. 376, 1947.
6. K. T. Dolder and O. Klemperer, High Frequency Oscillations in the Space Charge of some Electron Emission Systems, *Journal of Electronics*, **1**, p. 601 May, 1956.
7. C. N. Smyth, Total Emission Damping with Space-Charge-Limited Cathodes, *Nature*, **157**, p. 841, June 22, 1946.
8. W. Veith, Electron Energy Distribution in Space-Charge-Limited Electron Streams, *Zeit. f. angew. Physik*, **7**, No. 9, p. 437, 1955.
9. W. W. Rigrod and J. A. Lewis, Wave Propagation Along a Magnetically-Focused Cylindrical Electron Beam, *B.S.T.J.*, **33**, p. 399, March, 1954.
10. R. W. Peter, S. Bloom, and J. A. Ruetz, Space-Charge-Wave Amplification Along an Electron Beam by Periodic Change of the Beam Impedance, *RCA Rev.*, **15**, p. 113, March, 1954.
11. J. R. Pierce, The Wave Picture of Microwave Tubes, *B.S.T.J.*, **33**, p. 1343, Nov., 1954.
12. L. J. Chu, 1951 I.R.E. Electron Tube Conference on Electron Devices.
13. H. A. Haus and D. L. Bobroff, Small Signal Power Theorem for Electron Beams (to be published).
14. K. T. Dolder and O. Klemperer, Space-Charge Effects in Electron Optical Systems, *J. App. Phys.*, **26**, p. 1461, Dec., 1955.
15. S. J. Buchsbaum and E. Gordon, Highly Ionized Microwave Plasma, *M.I.T. R.L.E. Quarterly Prog. Rep.*, p. 11, Oct. 15, 1956.

Distortion Produced in a Noise Modulated FM Signal by Nonlinear Attenuation and Phase Shift

By S. O. Rice

(Manuscript received December 6, 1956)

An expression is given for the FM distortion introduced by a transducer whose attenuation and phase shift depend upon the frequency in an arbitrary way. This expression appears to be difficult to evaluate, but it yields useful approximations for the second and third order modulation terms. In all of the work, it is assumed that the distortion is small compared to the signal, and that the signal can be represented by a random noise having the same power spectrum.

INTRODUCTION

A number of workers have been concerned with the problem of computing the distortion introduced by a transducer when an FM wave passes through it. Some of the earliest results were published by Carson and Fry¹ and by van der Pol.² Several contributions to the subject have been made recently in connection with studies of microwave radio systems.

An excellent paper on this subject has been published recently by R. G. Medhurst and G. F. Small.³ Although their results differ considerably in form from those given here, they are nevertheless closely related to ours — their “sinusoidal variations of transmission characteristics” being special cases of our “nonlinear attenuation and phase shift.”

Here we treat the problem by applying a method used in a recent paper⁴ to study the distortion produced by an echo. Two assumptions are made, (1) that the distortion is small compared to the signal, and (2) that the signal can be represented by a random noise which has the same power spectrum as the signal. In Section I, we review some known results and put them in a form suited to our needs. Sections II and III are devoted to the derivation of our main formulas. The principal result is given by the triple integral (3.2) for the power spectrum of the dis-

tortion. Unfortunately, the integrals are difficult to evaluate. However, it is possible to obtain approximations for the second and third order modulation terms. These are given in Section IV. Some miscellaneous comments are made in Section V.

I APPROXIMATE EXPRESSION FOR THE DISTORTION $\theta(t)$

Let the FM signal be $\varphi'(t) = d\varphi/dt$ (for phase modulation the signal would be $\varphi(t)$). Then the FM wave is the real part of

$$v_i(t) = e^{ip t + i\varphi(t)} \quad (1.1)$$

where $p = 2\pi f_p$ is the carrier frequency. Let this wave pass through a transducer having attenuation α and phase shift β , where α and β are even and odd functions, respectively, of the frequency f . When a unit impulse of voltage $\delta(t)$ is applied to the transducer input, the output is

$$g(t) = \int_{-\infty}^{\infty} e^{-\alpha - i\beta + 2\pi i f t} df. \quad (1.2)$$

For physical systems, $g(t)$ is zero for negative t .

When $v_i(t)$ is applied to the transducer input, the output is

$$v_0(t) = \int_{-\infty}^{\infty} v_i(t') g(t - t') dt'. \quad (1.3)$$

When $v_0(t)$ is applied to an FM receiver, the detector output consists of the original signal $\varphi'(t)$ plus the distortion $\theta'(t)$ introduced by the transducer. Comparison with (1.1) shows that $\theta(t)$ may be obtained by solving

$$V(t) e^{ip t + i\varphi(t) + i\theta(t)} = v_0(t) \quad (1.4)$$

when p , $\varphi(t)$, $v_0(t)$ are assumed to be known, and $V(t)$, $\theta(t)$ unknown. When $V(t)$ is taken to be positive, (1.4) determines $\theta(t)$ except for an additive term of $2\pi n$ where n is an integer.

We now assume that the transducer acts like a good transmission medium in that the output differs but little from the input. More precisely, we assume

$$|v_0(t) - v_i(t)| \ll 1. \quad (1.5)$$

Since $|v_i(t)| = 1$, it follows that $|v_0(t)| \approx 1$. Transducers having appreciable attenuation and delay may be regarded as two transducers in tandem, one with constant (independent of f) values of α and β/f which are roughly equal to those of the original transducer, and the second

with variable α and β/f . The first transducer produces no distortion of the signal, and if condition (1.5) is satisfied by the second, the considerations of this paper will apply.

Equation (1.4) may be written as

$$V(t)e^{i\theta(t)} = v_0(t)/v_i(t)$$

so that

$$\theta(t) = \text{Im} \log \frac{v_0(t)}{v_i(t)}. \quad (1.6)$$

When we write

$$v_0(t)/v_i(t) = 1 + [v_0(t) - v_i(t)]/v_i(t),$$

expand the logarithm in (1.6), and use (1.5), we obtain our approximate expression for $\theta(t)$:

$$\begin{aligned} \theta(t) &= \text{Im} [v_0(t) - v_i(t)]/v_i(t) = \text{Im} v_0(t)/v_i(t) \\ &= \text{Im} [v_i(t)]^{-1} \int_{-\infty}^{\infty} v_i(t') g(t - t') dt' \\ &= \text{Im} \int_{-\infty}^{\infty} \exp[ip(t' - t) + i\varphi(t') - i\varphi(t)] g(t - t') dt'. \end{aligned} \quad (1.7)$$

So far there is nothing essentially new in our work.⁵

II AUTOCORRELATION FUNCTION OF $\theta(t)$

In Section I, $\varphi'(t)$ could be any reasonable sort of signal. In the following work we assume that it is a Gaussian noise whose power spectrum, $w_{\varphi'}(f)$, is given to us. The power spectrum of $\varphi(t)$ is

$$w_{\varphi}(f) = w_{\varphi'}(f)/(2\pi f)^2, \quad (2.1)$$

and its autocorrelation function is

$$\psi_{\tau} = \int_0^{\infty} w_{\varphi}(f) \cos 2\pi f\tau df. \quad (2.2)$$

We have written ψ_{τ} instead of $\psi(\tau)$ or $R_{\varphi}(\tau)$ to simplify the appearance of the formulas which occur in our work.

Our problem is to find the power spectrum, $w_{\theta}(f)$, of the distortion $\theta(t)$, given $w_{\varphi}(f)$. The method of solution is much the same as that used in Reference 4. We first find the autocorrelation function $R_{\theta}(\tau)$ of $\theta(t)$ and then obtain $w_{\theta}(f)$ by taking the Fourier cosine transform of $R_{\theta}(\tau)$.

Let the last integral in (1.7) be $F(t)$ so that $\theta(t) = \text{Im } F(t)$. Then

$$\theta(t)\theta(t + \tau) = \frac{1}{2} \text{Re} \{F(t)F^*(t + \tau) - F(t)F(t + \tau)\} \quad (2.3)$$

where $F^*(t + \tau)$ is the complex conjugate of $F(t + \tau)$. The autocorrelation function of $\theta(t)$ is obtained by averaging over the ensemble of the noise functions $\varphi(t)$:

$$\begin{aligned} R_\theta(\tau) &= \text{av } \theta(t)\theta(t + \tau) \\ &= \text{av } \frac{1}{2} \text{Re} \left\{ \int_{-\infty}^{\infty} dt' \int_{-\infty}^{\infty} dt'' \exp [ip(t' - t) + i\varphi(t') - i\varphi(t)] \right. \\ &\quad \cdot g(t - t')g(t + \tau - t'') [\exp [-ip(t'' - t - \tau) \\ &\quad - i\varphi(t'') + i\varphi(t + \tau)] - \exp [ip(t'' - t - \tau) + i\varphi(t'') \\ &\quad \left. - i\varphi(t + \tau)] \right\}. \end{aligned} \quad (2.4)$$

Since $g(t)$ is real, $g^*(t) = g(t)$. The averaging process may be carried out by a method analogous to that used in Reference 4. The formula to be used is

$$\begin{aligned} \text{av } \exp [i\varphi(t') - i\varphi(t) + ia\varphi(t'') - ia\varphi(t + \tau)] \\ = \exp [-\psi_0(1 + a^2) + \psi_{t'-t} - a\psi_{t'-t''} + a\psi_{t'-t-\tau} \\ + a\psi_{t-t''} - a\psi_\tau + a^2\psi_{t'-t-\tau}] \end{aligned} \quad (2.5)$$

where a is either -1 or $+1$, and ψ_τ is an even function of τ . When (2.5) is used in (2.4) a double integral for $R_\theta(\tau)$ is obtained. The substitutions

$$\begin{aligned} x &= t - t', \\ y &= t + \tau - t'', \\ R_v &= \psi_{\tau+x-y} - \psi_{\tau+x} - \psi_{\tau+x} - \psi_{\tau-y} + \psi_\tau \end{aligned} \quad (2.6)$$

convert the double integral into

$$\begin{aligned} R_\theta(\tau) &= \frac{1}{2} \text{Re} \int_{-\infty}^{\infty} dx \int_{-\infty}^{\infty} dy g(x) e^{-ipx-2\psi_0+\psi_x+\psi_y} \\ &\quad \cdot g(y) [e^{ipy+R_v} - e^{-ipy-R_v}]. \end{aligned} \quad (2.7)$$

The symbol R_v is chosen to agree as closely as possible with the notation of Reference 4. There R_v was the autocorrelation of the random function, $v(t)$, where $v(t + T) = \varphi(t) - \varphi(t + T)$, T being the echo delay. Here, R_v is the average value of the product,

$$[\varphi(t) - \varphi(t + y)][\varphi(t + \tau) - \varphi(t + \tau + x)]$$

which becomes the autocorrelation function of $v(t)$ when $y = x = T$.

It may be verified that the expression (2.7) for $R_\theta(\tau)$ is an even function of τ , as it should be. Expression (2.7) is the autocorrelation function we set out to find.

The distortion $\theta(t)$ has an average value, $\bar{\theta}$, whose square is $R_\theta(\infty)$. Since $\varphi(t)$ is a noise function, its autocorrelation function ψ_τ goes to zero as τ approaches ∞ . Hence, $R_\theta(\infty)$ is given by the expression obtained from (2.7) by setting $R_v = 0$. The autocorrelation function of $\theta(t) - \bar{\theta}$ is

$$\begin{aligned} R_{\theta-\bar{\theta}}(\tau) &= R_\theta(\tau) - R_\theta(\infty) \\ &= \frac{1}{2} \operatorname{Re} \int_{-\infty}^{\infty} dx \int_{-\infty}^{\infty} dy g(x) e^{-ipx-2\psi_0+\psi_x+\psi_y} \\ &\quad \cdot g(y) [e^{ipy}(e^{R_v} - 1) - e^{-ipy}(e^{-R_v} - 1)]. \end{aligned} \quad (2.8)$$

III POWER SPECTRUM OF THE DISTORTION

Since $\theta(t)$ has an average value which is generally not zero, its power spectrum, $w_\theta(f)$, has a spike of infinite height at $f = 0$ corresponding to the power in the dc component $\bar{\theta}$. When this spike is subtracted from $w_\theta(f)$ the remainder is the power spectrum of $\theta(t) - \bar{\theta}$ given by

$$w_{\theta-\bar{\theta}}(f) = 4 \int_0^{\infty} R_{\theta-\bar{\theta}}(\tau) \cos 2\pi f \tau d\tau. \quad (3.1)$$

When we use (2.8) and note that $R_{\theta-\bar{\theta}}(\tau)$ is an even function of τ , we obtain

$$\begin{aligned} w_{\theta-\bar{\theta}}(f) &= \int_{-\infty}^{\infty} dx \int_{-\infty}^{\infty} dy g(x) g(y) e^{-2\psi_0+\psi_x+\psi_y} \int_{-\infty}^{\infty} [\cos(px - py) \\ &\quad \cdot (e^{R_v} - 1) - \cos(px + py)(e^{-R_v} - 1)] \cos 2\pi f \tau d\tau. \end{aligned} \quad (3.2)$$

Reasoning similar to that given in Reference 4 shows that the inter-channel interference spectrum, $w_c(f)$, (i.e., $w_c(f)\Delta f$ is the average amount of distortion power received in an idle channel of width Δf centered on the frequency f , all other channels being busy) may be obtained from (3.2) by replacing $(e^{\pm R_v} - 1)$ by $(e^{\pm R_v} \mp R_v - 1)$.

The power spectrum of $\theta(t) - \bar{\theta}$ may be regarded as made up of modulation products of all orders. It turns out that the contribution of n^{th} order products is given by the integral of the R_v^n terms obtained from the power series expansions of $\exp[\pm R_v]$.

IV FIRST AND SECOND ORDER MODULATION TERMS

Here we shall study the first and second order modulation terms. These arise from the first and second powers of R_v in the expansion of

the quantity within the square brackets in (3.2):

$$2R_v \cos px \cos py + R_v^2 \sin px \sin py. \quad (4.1)$$

The integrations with respect to τ may be performed with the help of

$$\int_{-\infty}^{\infty} \psi_{\tau+b} \cos 2\pi f \tau \, d\tau = \frac{w_{\varphi}(f)}{2} \operatorname{Re} e^{-i2\pi fb}, \quad (4.2)$$

$$\int_{-\infty}^{\infty} \psi_{\tau+b} \psi_{\tau+c} \cos 2\pi f \tau \, d\tau \quad (4.3)$$

$$= \operatorname{Re} \frac{1}{4} \int_{-\infty}^{\infty} du \, w_{\varphi}(u) w_{\varphi}(f-u) \exp \{-i2\pi[bu + c(f-u)]\}$$

which follow from (2.2) and the fact that we have defined $w(-f)$ to be equal to $w(f)$. In our notation the total power in a random noise function is the integral of $w(f)$ from $f = 0$ to $f = \infty$.

The first order modulation term is obtained from (3.2) by replacing the term within the square bracket by $2R_v \cos px \cos py$. When the expression (2.6) for R_v is used, the integration with respect to τ may be performed with the help of (4.2):

$$\int_{-\infty}^{\infty} R_v \cos 2\pi f \tau \, d\tau = \frac{w_{\varphi}(f)}{2} \operatorname{Re} [(e^{-2\pi i x f} - 1)(e^{i2\pi y f} - 1)]. \quad (4.4)$$

This leads to the following expression for the first order modulation term in (3.2)

$$w_{\varphi}(f) \left| \int_{-\infty}^{\infty} dx g(x) e^{-\psi_0 + \psi_x} \cos px (e^{-2\pi i x f} - 1) \right|^2. \quad (4.5)$$

This is the quantity which is to be subtracted from $w_{\theta-\bar{\theta}}(f)$ to obtain the interchannel interference spectrum $w_c(f)$.

The second order modulation term is handled in much the same manner. With the help of (4.3) it may be shown that

$$\begin{aligned} \int_{-\infty}^{\infty} R_v^2 \cos 2\pi f \tau \, d\tau &= \operatorname{Re} \frac{1}{4} \int_{-\infty}^{\infty} du \, w_{\varphi}(u) w_{\varphi}(f-u) \\ &\cdot (e^{-2\pi i x u} - 1) (e^{-2\pi i x (f-u)} - 1) \\ &\cdot (e^{2\pi i y u} - 1) (e^{2\pi i y (f-u)} - 1). \end{aligned} \quad (4.6)$$

From this it follows that the second order modulation term in (3.2) is

$$\begin{aligned} \frac{1}{2! 2} \int_{-\infty}^{\infty} du \, w_{\varphi}(u) w_{\varphi}(f-u) &\left| \int_{-\infty}^{\infty} dx g(x) e^{-\psi_0 + \psi_x} \sin px \right. \\ &\cdot (e^{-2\pi i x u} - 1) (e^{-2\pi i x (f-u)} - 1) \left. \right|^2. \end{aligned} \quad (4.7)$$

When $\psi_0 - \psi_x$ is so small that $\exp(-\psi_0 + \psi_x)$ may be replaced by unity, as it is in some important practical cases, approximations may be obtained for (4.5) and (4.7). The integral in x may be expressed as the sum of integrals of the type

$$\int_{-\infty}^{\infty} g(x) e^{-ipx-2\pi iax} dx = [e^{-\alpha-i\beta}]_{f=a+f_p} \\ = G_a + iB_a, \quad (4.8)$$

$$\int_{-\infty}^{\infty} g(x) e^{ipx-2\pi iax} dx = G_{-a} - iB_{-a}.$$

The values of the integrals follow from (1.2) and the Fourier integral theorem. G and B are, respectively, even and odd functions of frequency, and G_a , B_a are their values at the frequency $f = f_p + a$ where $f_p = p/2\pi$ is the carrier frequency:

$$G \text{ at frequency } f_p + a = G_a,$$

$$B \text{ at frequency } f_p + a = B_a.$$

In this way we get the approximation

$$4^{-1} w_{\varphi}(f) [(G_f - 2G_0 + G_{-f})^2 + (B_f - B_{-f})^2] \quad (4.9)$$

for the first order modulation term, and

$$\frac{1}{2!8} \int_{-\infty}^{\infty} du w_{\varphi}(u) w_{\varphi}(f-u) [(G_u - G_{-u} + G_{f-u} - G_{-f+u} - G_f \\ + G_{-f})^2 + (B_u + B_{-u} + B_{f-u} + B_{-f+u} - B_f - B_{-f} - 2B_0)^2] \quad (4.10)$$

for the second order modulation term.

Expression (4.10) is an approximation to the second order modulation term (4.7). When most of the interchannel interference is due to second order modulation products, (4.10) is also an approximation to $w_c(f)$, the interchannel interference spectrum. The following remarks may be of some help in deciding whether (4.10) may be used.

1. For the case of phase modulation and a "flat" signal band, the first of equations (5.3) shows that ψ_0 and ψ_r may be made as small as we please by choosing the signal power (as measured by P_0) small enough. Since R_v is proportional to P_0 , P_0 may be chosen small enough to make R_v^3 and higher order terms negligible in the expansion of the integrand of (3.2) (unless there is some sort of symmetry which causes the second order terms to vanish). In this case the interference is mostly second order modulation and (4.7) is a good approximation to $w_c(f)$. Furthermore, as P_0 approaches zero, $\exp(-\psi_0 + \psi_x)$ approaches unity

and (4.10) becomes a good approximation to (4.7). Just how small P_0 has to be depends upon the signal bandwidth, f_b , and the characteristics of the transducer.

2. For the case of FM and a flat signal band, the second of equations (5.3) shows that even if P_0 is small, the difference $\psi_0 - \psi_\tau$ approaches ∞ as $|\tau|$ approaches ∞ . To justify the use of (4.10) in this case it is necessary to take into account the behavior of $g(t)$, the response of the transducer to the unit impulse $\delta(t)$. For example, if the duration of $g(x)$ in (4.7) is so brief that $g(x)$ becomes negligibly small before $-\psi_0 + \psi_x$ becomes appreciably different from zero (which may be achieved by making P_0 small enough) then (4.10) is a good approximation to (4.7).

3. When the attenuation, α , and phase shift, β , are given for any particular transducer, the corresponding $g(t)$ may be obtained from (1.2). Once $g(t)$ and $\psi_0 - \psi_\tau$ are known, the conditions under which $\exp(-\psi_0 + \psi_x)$ may be replaced by unity in (4.7) and $O(R_v^3)$ terms neglected in (3.2) may be determined by direct examination of the integrals.

As might be expected, the third order modulation results are quite complicated. The third order modulation term in (3.2) is

$$\frac{1}{3!4} \int_{-\infty}^{\infty} df' \int_{-\infty}^{\infty} df'' w_\varphi(f') w_\varphi(f'') w_\varphi(f''') \left| \int_{-\infty}^{\infty} dx g(x) \cos px e^{-\psi_0 + \psi_x} (z^{f'} - 1)(z^{f''} - 1)(z^{f'''} - 1) \right|^2 \quad (4.11)$$

where $f''' = f - f' - f''$ and $z = \exp(-i2\pi x)$. When ψ_0 is small this is approximately

$$\frac{1}{3!16} \int_{-\infty}^{\infty} df' \int_{-\infty}^{\infty} df'' w_\varphi(f') w_\varphi(f'') w_\varphi(f''') [H^2 + K^2] \quad (4.12)$$

where

$$H = m(f') + m(f'') + m(f) + m(f - f' - f'') - m(f - f') - m(f - f'') - m(0) - m(f' + f''), \quad (4.13)$$

$$m(f) = G_f + G_{-f}, \quad n(f) = B_f - B_{-f},$$

and K is an expression obtained from H by replacing n by m .

V MISCELLANEOUS COMMENTS

Here we make some miscellaneous comments related to the foregoing results.

If the transducer is perfect except for an echo, its response to a unit impulse $\delta(t)$ is

$$g(t) = \delta(t) + r\delta(t - T) \quad (5.1)$$

where r and T are the amplitude and the delay of the echo. The results obtained using (5.1) agree, as they should, with the results obtained in Reference 4. Of course, r must be assumed small compared to unity in order that condition (1.5) may hold.

When the power spectrum of the signal is equal to a constant P_0 over the band (f_a, f_b) and zero elsewhere we have for phase and frequency modulation, respectively,

$$\begin{aligned} \text{PM: } w_\varphi(f) &= P_0, & f_a < f < f_b, \\ \text{FM: } w_\varphi(f) &= P_0/(2\pi f)^2, & f_a < f < f_b. \end{aligned} \quad (5.2)$$

When $f_a = 0$ the autocorrelation functions are

$$\begin{aligned} \text{PM: } \psi_\tau &= P_0 f_b (\sin v)/v, \\ \text{FM: } \psi_0 - \psi_\tau &= A[-1 + \cos v + v \text{Si}(v)], \\ v &= 2\pi f_b \tau, \quad A = P_0 f_b (2\pi f_b)^{-2} = (\sigma/f_b)^2. \end{aligned} \quad (5.3)$$

The mean square values of the signals are $P_0 f_b$ (radians)² for PM and $P_0 f_b$ (radians/sec)² for FM. If, for FM, σ is the rms frequency deviation in cps (so that the "peak" deviation is, say, 4σ cps) then $(2\pi\sigma)^2 = P_0 f_b$. The difference $\psi_0 - \psi_\tau$ is used in the FM case to avoid difficulty at $f = 0$. It will be noticed that our formulas are such that the ψ 's may be replaced by $(\psi - \psi_0)$'s without altering the values of the various exponents, etc. In microwave systems the quantity A is often small in comparison with unity.

As an example of the use of the second order modulation approximation (4.10) consider the case where the attenuation, α , is zero and the phase shift $\beta = a_2(f - f_p)^2/2$ radians, a_2 being small. Then, since $G \approx 1 - \alpha$, $\beta \approx -\beta$, we have $G_u \approx 1$ and

$$\begin{aligned} B_u &\approx -[\beta \text{ for } f = f_p + u] \\ &= -a_2 u^2/2. \end{aligned} \quad (5.4)$$

When we take the FM case of (5.2) and substitute in the approximation (4.10), the interchannel interference power spectrum is found to be

$$\begin{aligned} \frac{1}{2!8} \int_{f-f_b}^{f_b} \frac{P_0}{(2\pi u)^2} \frac{P_0}{(2\pi)^2(f-u)^2} [0 + (2a_2 u(f-u))^2] du \\ = (2\pi)^{-4} (a_2 P_0/2)^2 (2f_b - f). \end{aligned} \quad (5.5)$$

Dividing by $w_{\varphi}(f) = P_0/(2\pi f)^2$ gives the ratio of the interference power to the signal power

$$(a_2\sigma f/2)^2(2 - f/f_b) \quad (5.6)$$

where the relation $P_0 = (2\pi\sigma)^2/f_b$ has been used to eliminate P_0 . Here σ is the rms frequency deviation of the FM signal in cps. The expression (5.6) agrees with results of some earlier work done at Bell Telephone Laboratories. In that work the second order modulation products were summed directly.

It is interesting to apply the formulas given here to some of the cases considered by Medhurst and Small.³ They have shown that when (in our notation) $\alpha = -r \cos 2\pi fT$ and $\beta = 0$ the power spectrum of the distortion is

$$w_{\theta-\bar{\theta}}(f) = \sin^2 \pi fT [w_{\theta-\bar{\theta}}(f)]_{\text{echo}}, \quad (5.7)$$

and when $\alpha = 0$ and $\beta = r \sin 2\pi fT$,

$$w_{\theta-\hat{\theta}}(f) = \cos^2 \pi fT [w_{\theta-\hat{\theta}}(f)]_{\text{echo}}. \quad (5.8)$$

Here $[w_{\theta-\bar{\theta}}(f)]_{\text{echo}}$ is the power spectrum of the distortion due to a simple echo of amplitude r and delay T (corresponding to $\alpha = -r \cos 2\pi fT$ and $\beta = r \sin 2\pi fT$). Expressions (5.7) and (5.8) may also be obtained by setting the impulse response $g(t)$ equal to

$$\delta(t) + \frac{r}{2} \delta(t - T) \pm \frac{r}{2} \delta(t + T)$$

in (3.2).

The second order modulation approximation for the $\alpha = -r \cos 2\pi fT$, $\beta = 0$ case may be obtained from (4.10) and turns out to be

$$\int_{-\infty}^{\infty} w_{\varphi}(u)w_{\varphi}(f-u)[2r \sin pT \sin \pi fT \sin \pi uT \sin \pi(f-u)T]^2 du. \quad (5.9)$$

It is seen that this contains the factor $\sin^2 \pi fT$ predicted by (5.7). When (5.9) is applied to the FM case of (5.2) an integral something like (5.5) (but more complicated) is obtained. The ratio of the second order modulation interference power to the signal power is found to be

$$2[r \sin pT \sin \pi fT]^2(\sigma/f_b)^2 UK \quad (5.10)$$

where K is the quantity

$$K = 2\alpha^2 \int_{\alpha-r}^U \left[\frac{\sin(y/2) \sin(\alpha-y)/2}{y(\alpha-y)} \right]^2 dy \quad (5.11)$$

tabulated in Table 4.2 of Reference 4 and

$$\alpha = 2\pi fT, \quad U = 2\pi f_b T. \quad (5.12)$$

The parameters a and k that appear in the table are defined by

$$a = f/f_b \quad \text{and} \quad k = 8f_b T.$$

These formulas serve to supplement the formulas and curves given by Medhurst and Small.

ACKNOWLEDGMENT

I wish to express my thanks to H. E. Curtis who has furnished some of the examples given in this paper and to E. D. Sunde, S. Doba, and others for their helpful comments.

REFERENCES

1. J. R. Carson and T. C. Fry, Variable Frequency Electric Circuit Theory With Application to the Theory of Frequency Modulation, B.S.T.J., **16**, 510-540, Oct., 1937.
2. B. van der Pol, The Fundamental Principles of Frequency Modulation, J.I.E.E., **93**, pp. 153-158, 1946.
3. R. G. Medhurst and G. F. Small, Distortion in Frequency-Modulation Systems Due to Small Sinusoidal Variations of Transmission Characteristics, Proc. I.R.E., **44**, pp. 1608-1612, Nov., 1956.
4. W. R. Bennett, H. E. Curtis, and S. O. Rice, Interchannel Interference in FM and PM Systems Under Noise Loading Conditions, B.S.T.J., **34**, pp. 601-636, May, 1955. The same problem has been treated independently and in much the same way by S. V. Borodich, On the Nonlinear Distortions Caused by Variations of the Antenna Feeder in Multichannel Frequency Modulation Systems, Radiotekhnika, Moscow, **10**, pp. 3-14, 1955.
5. These results are closely associated with some given by M. K. Zinn, Transient Response of an FM Receiver, B.S.T.J., **29**, pp. 714-731, 1948.

Self-Timing Regenerative Repeaters

By E. D. SUNDE

(Manuscript received March 29, 1956)

In self-timing regenerative repeaters, a timing wave for control in pulse regeneration is derived from the binary pulse train at each repeater with the aid of a resonant circuit tuned to the pulse repetition frequency. The timing wave can be made to exercise complete control in retiming of pulses independent of the received pulse train, or it can be combined with the received pulse train to provide partial retiming. The timing principles are discussed here for a particular type of self-timed regenerative repeater invented by Wrathall, in which a timing wave derived from either the received or the regenerated pulse train is combined in a particular way with the received pulse train. The regeneration characteristics of such repeaters as determined by various design parameters are investigated, together with the cumulation of timing deviations in repeater chains and the circuit requirements that must be met to insure satisfactory performance.

INTRODUCTION

Pulse transmission systems employing binary codes, such as PCM, have two inherent properties that are desirable from the standpoint of avoiding excessive transmission impairments by noise and other imperfections in the transmission medium. For one thing binary pulse codes permit substantial transmission distortion of pulses within certain tolerable limits with negligible degradation of received signals. For another, regenerative repeaters can be used at intervals along a route to prevent accumulation of transmission distortion of pulses from various sources, so that virtually the entire allowable distortion can be permitted in each link or repeater section.

The above desirable properties are secured in exchange for increased channel bandwidth, and can be used to full advantage in applications of binary pulse systems to such transmission media as radio and wave guides, where transmission is at such high frequencies that increased channel bandwidth does not entail increased attenuation. In wire circuits, however, where baseband transmission is the more economical

method, attenuation increases nearly in proportion to the square root of the channel bandwidth. For this reason, rather short repeater spacings may be required for binary pulse systems, so that for economical applications to wire circuits it is imperative to have reliable regenerative repeaters of simple design.

In their principle of operation regenerative repeaters are by nature more complicated than ordinary repeaters. In addition to providing gain to off-set attenuation in the transmission medium, as in ordinary repeaters, they must also perform gating operations for sampling and regenerating the received pulse train. This, however, does not preclude the possibility that these operational principles can be implemented in repeater design by instrumentation that is simpler than required for ordinary repeaters.

The possibility of simple instrumentation resides partly in the circumstance that equalization circuitry for regenerative repeaters can be substantially simpler than for ordinary repeaters, owing to less exacting requirements on equalization. Furthermore, satisfactory performance in pulse regeneration can be achieved without very precise timing in sampling and regeneration of pulse trains. It is thus possible to secure nearly the same performance as for ideal regenerative repeaters by partial rather than complete exact retiming of pulse trains at each repeater. This facilitates simple gating arrangements for regeneration of pulses. Moreover, it permits a timing wave for control of gating operations to be derived from either the received or regenerating pulse trains with the aid of a simple resonant of circuit.

The simplicity of instrumentation permitted by these considerations is exemplified in a self-timed regenerative repeater for baseband pulses invented by L. R. Wrathall of Bell Telephone Laboratories. The circuitry of the repeater together with the results of tests on laboratory models are dealt with elsewhere¹ and not considered here. The purpose of this paper is an analysis of the timing principles underlying this type of repeater together with its regeneration characteristics as determined by various basic design parameters, on the assumption of ideal implementation of the timing principles by appropriate instrumentation. In the Wrathall repeater "quantized feed-back" is employed as a means of reducing the effect of low-frequency cut-off in transformers. Since this is not an essential feature of self-timing repeaters and has no direct bearing on the timing principles, it is disregarded herein.

¹ L. R. Wrathall, Transistorized Binary Pulse Regenerator, B.S.T.J., **35**, pp. 1059-1084, Sept., 1956.

I REGENERATION AND RETIMING

1.0 General

In an ideal regenerative repeater the received pulse train is sampled at proper fixed intervals, to determine whether a pulse is present. The regenerated pulses transmitted into the next repeater section are all of the same shape and amplitude, independent of the shape of the input pulses. Thus pulse distortion from noise and other system imperfections is removed, provided the maximum distortion is held within proper limits. Errors in the form of pulses in place of spaces, or conversely, are encountered when these limits are exceeded. In a repeater chain there will be cumulation of errors in proportion to the number of repeater sections in tandem. However, the rate of errors in each section and thus in the whole chain can be limited by a relatively small increase in the signal-to-noise ratio of each section as the number of repeaters in tandem is increased. This increase in signal-to-noise ratio with increasing length of the repeater chain is much less than with ordinary nonregenerative repeaters. For this reason regenerative rather than ordinary repeaters are desirable, though not essential for systems employing binary codes.

An ideal regenerative repeater with the above features would entail rather complicated instrumentation for precise timing, sampling and pulse regeneration. With partial rather than complete exact retiming the repeaters can be simplified, in exchange for some sacrifice in performance, as shown later.

1.1 Regeneration Without Retiming

It would be possible to have a repeater in which pulses would be regenerated in amplitude and shape, but without retiming. Pulses would in this case be regenerated when the amplitude of the pulse exceeded a certain triggering level L . If the pulse shape is given by $P(t)$, this would occur at a time t_0 such that

$$P(t_0) = L. \quad (1.1)$$

This would permit simple instrumentation, since regenerated pulses would be triggered without separate sampling of the received pulse train. With this method, however, timing deviations in the regenerated pulses would result from transmission distortion of the received pulses by noise and other system imperfections. These timing deviations would cumulate in a repeater chain and cause a reduction in the tolerance of the repeaters to noise, such that the signal-to-noise ratio would have to

be increased with the number of repeaters in tandem in the same way as for ordinary repeaters.

1.2 Regeneration with Complete Retiming

With complete retiming, the instants of pulse regeneration would be controlled by a periodic retiming wave, $R(t)$, with a fundamental period equal to the interval between pulses. The received pulse train would be sampled at instants when the retiming wave had a certain level L_s . The sampling instants t_0 would thus be given by

$$R(t_0) = L_s. \quad (1.2)$$

$R(t_0)$ would satisfy this equation for $t_0 = nT \pm \Delta T$, where T is the nominal interval between pulses, n is an integer and ΔT is a certain tolerable deviation from the desired sampling instants. Pulses would be regenerated provided $P(t_0) > L$ and would be omitted if $P(t_0) < L$.

With this method the timing deviations in regenerated pulses would be limited to $\pm \Delta T$, regardless of the timing deviations in received pulses. There would be no cumulation of timing deviations in a repeater chain. However, the tolerance of the repeaters to noise would be somewhat reduced by the timing deviations $\pm \Delta T$.

1.3 Regeneration with Partial Retiming

Partial retiming is obtained by a combination of the above two methods, by triggering regenerated pulses without sampling at instants t_0 determined by

$$P(t_0) + R(t_0) = L. \quad (1.3)$$

To permit regeneration without sampling and without a marked reduction in the tolerance of the repeaters to noise, the timing wave $R(t)$ must meet certain conditions illustrated in Fig. 1. One is that it must be a nearly periodic function as for complete retiming. The second condition is that $R(t)$ must be zero near the sampling points to obtain substantially the same tolerance to noise in the presence of a pulse as in the absence of a pulse. A third condition is that $R(t)$ must have substantial negative values between sampling points in order that the repeater be rather insensitive to noise between sampling points, as with complete retiming. It will be recognized that, in general, the maximum value of $R(t)$ need not necessarily be zero, as in the above illustration. It can be greater or smaller than zero, provided the triggering level is

modified accordingly. A maximum value of zero is, however, convenient from the standpoint of instrumentation.

A limiting shape of retiming wave that would result in complete re-timing, but without the need for special sampling is also illustrated in Fig. 1.

1.4 Derivation of Timing Wave from Pulse Train

As shown above, the retiming wave must be essentially periodic, with a fundamental frequency equal to the pulse repetition frequency $f = 1/T$, where T is the interval between pulses. The simplest form is a sinusoidal wave, which can be derived from the pulse train at repeaters with the aid of a narrow band-pass filter, such as a simple resonant circuit cen-

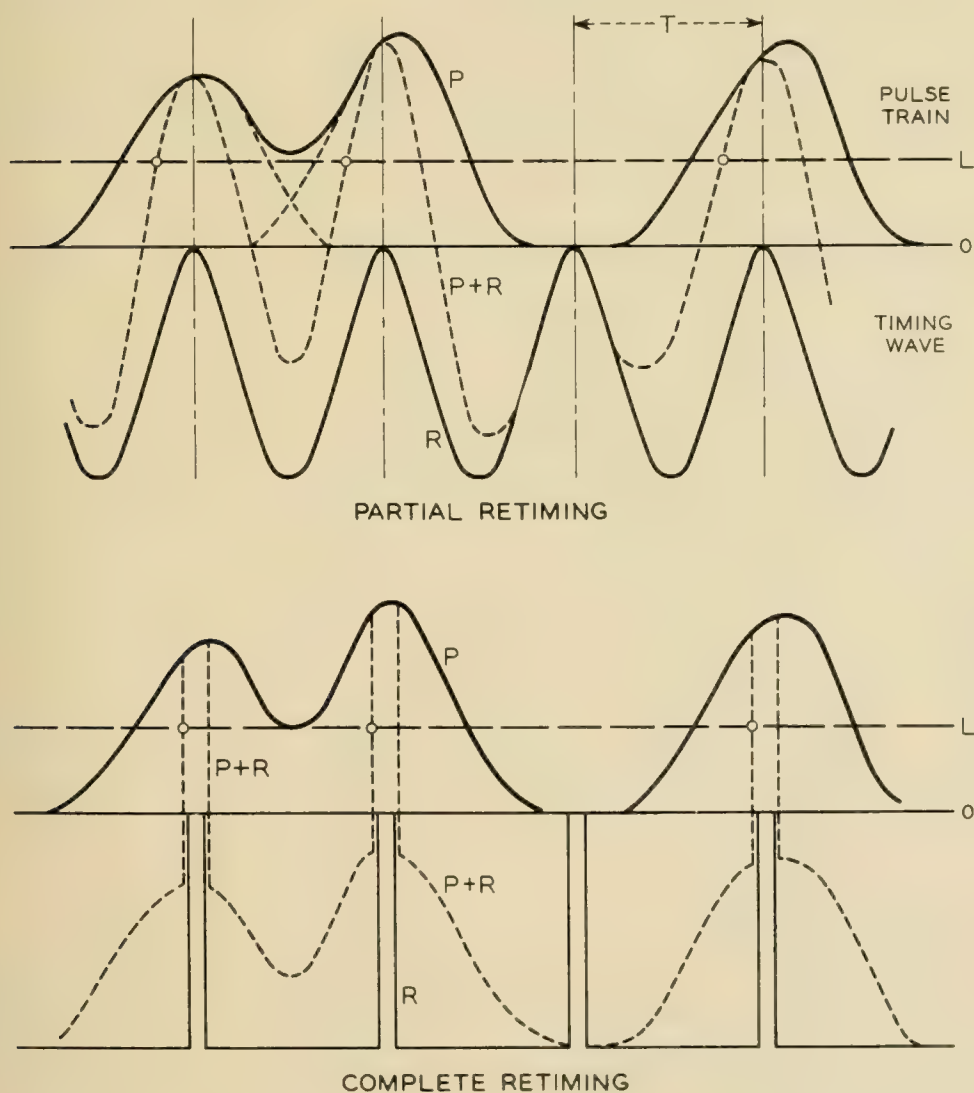


Fig. 1 — Principle of partial retiming method.

tered on the pulse repetition frequency. This possibility resides in the circumstance that a random "on-off" pulse train can be resolved into two components. One is an infinite sequence of pulses of the same polarity and equal amplitude, the other a sequence of randomly positive and negative polarity. The response of a resonant circuit to the first component is a steady state sinusoidal wave of the pulse repetition frequency. The second component gives rise to random variations in amplitude and phase, which in principle can be limited to any desired extent by limiting the band of the resonant circuit and the deviation in the resonant frequency from the pulse repetition frequency.

A principal feature of this method of "self-timing", aside from its simplicity, is that the timing wave becomes a slave of the pulse train. Thus, if there is a fixed delay in pulse regeneration at a repeater, the same delay is imparted to the timing wave derived from the pulse train at the next repeater. This prevents a cumulation of such fixed delays with respect to the timing wave, but not with respect to an absolute time scale; i.e., with respect to an ideal timing wave transmitted along the repeater chain and independent of the pulse train.

1.5 Self-Timed Repeaters with Partial Retiming

A timing wave derived from the pulse train with the aid of a resonant circuit can be used in conjunction with complete or partial retiming. With complete retiming, pulses could be regenerated at the zero points in the timing wave, and the effects of amplitude variations in the timing wave can thus be avoided. Timing deviations in the regenerated pulses would in this case depend only on phase deviations in the timing wave, caused partly by the component of randomly positive and negative polarity in the pulse train and partly by timing deviations in the pulse train from which the timing wave is derived.

With partial retiming the situation is more complex. Timing deviations in regenerated pulses in this case depend not only on amplitude and phase variations in the timing wave, but also on the regeneration characteristics of the repeaters.

1.6 Types of Timing Deviations

In a regenerated pulse train there will be fixed and random timing deviations. Of the latter there are three types. One is the timing deviation taken in relation to an exact timing wave with a period T equal to the nominal pulse interval. The second is the timing deviation taken in relation to the timing wave derived from the pulse train, which in itself

will contain random deviations. The third type is random deviations in the interval of adjacent pulses. If the first type is held within tolerable limits, this will also be the case for the second and third types. For this reason only the first type is considered herein.

II REGENERATION CHARACTERISTICS WITH PARTIAL RETIMING

2.0 General

With partial retiming, there will be timing deviations in the regenerated pulses as a result of timing deviations, amplitude variations and distortion by noise of both the received pulses and the timing wave.

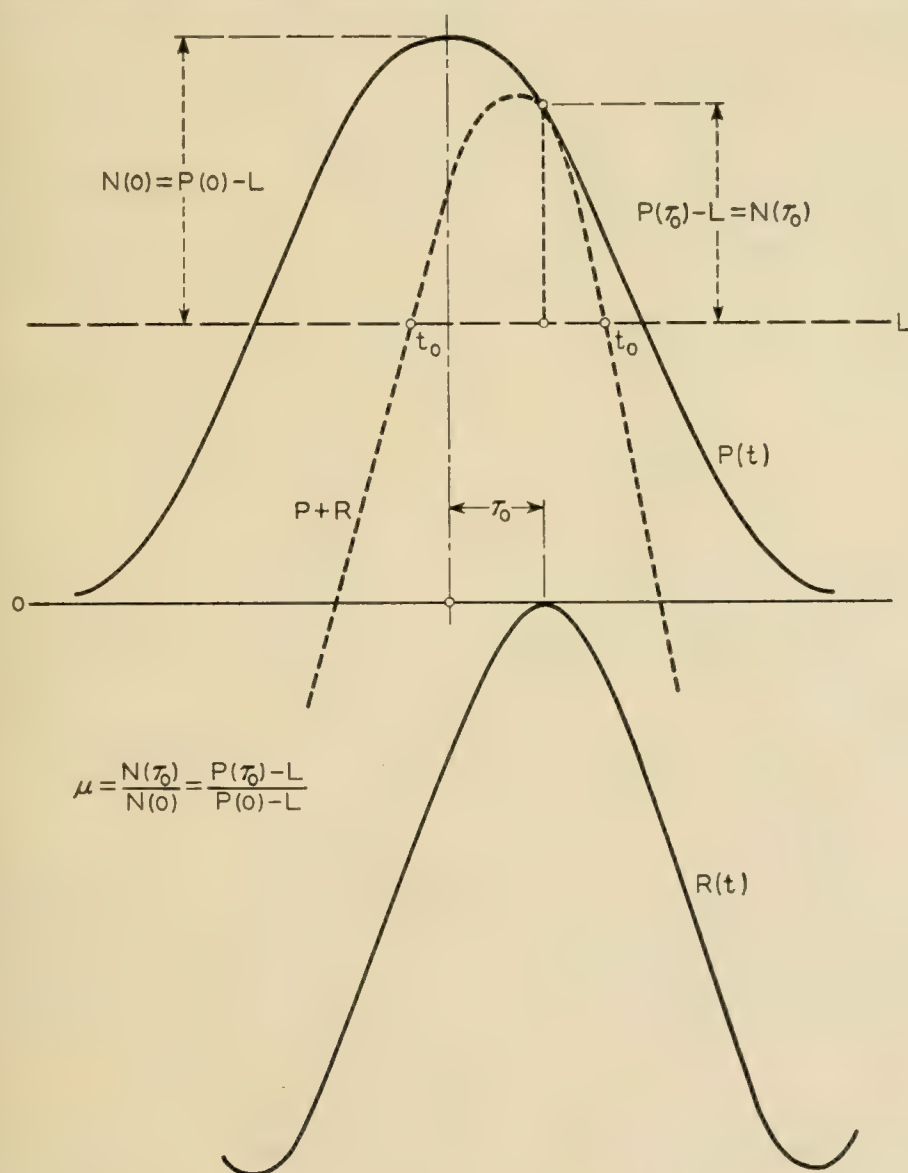


Fig. 2 — Reduction in tolerance to noise by displacement in timing wave.

The conversion of these variations into timing deviations in the regenerated pulses depends on certain relationships between the pulse train and the timing wave, discussed in the following sections.

2.1 Tolerance to Noise

From Fig. 2 it can be seen that if the timing wave is displaced by τ_0 , the value of $P(t) + R(t - \tau_0)$ in the presence of a pulse exceeds the triggering level by a maximum amount

$$[P(t) + R(t - \tau_0) - L]_{\max} \cong [P(\tau_0) - L]. \quad (2.1)$$

It will be recognized that the right-hand side of this equation represents the tolerance to noise of negative amplitudes with instantaneous sampling at $t = \tau_0$, as in an ideal repeater with complete retiming.

With partial retiming, the tolerance to noise will be less than the above maximum value. However, it will be greater than the average of $P(t) + R(t - \tau_0) - L$ in the range where the latter difference is positive. Let it be assumed that it is smaller than the maximum by a factor k somewhat smaller than unity. The tolerance to noise with a displacement τ_0 in the timing wave is then smaller than without a displacement (i.e., $\tau_0 = 0$) by the factor

$$\mu = \frac{k[P(\tau_0) - L]}{k[P(0) - L]} = \frac{P(\tau_0) - L}{P(0) - L}. \quad (2.2)$$

The tolerance to noise will thus be reduced in a way similar to that for an ideal repeater with complete retiming. The absolute tolerance to noise will be less than for a repeater with complete retiming by a factor k somewhat smaller than unity, say in the order 0.8, corresponding to about 2 db.

2.2 Conversion of Timing Deviations

With partial retiming, timing deviations in received pulses and in the timing wave are converted into smaller deviations in regenerated pulses.

Let τ_p be a time displacement in a received pulse and τ_r in the timing wave, both in the positive direction. Pulses will then be regenerated at a time t_0' given by

$$P(t_0' - \tau_p) + R(t_0' - \tau_r) = L \quad (2.3)$$

where the minus signs are used since this corresponds to a displacement of P and R in the positive direction. Subtracting (1.3) from (2.3),

$$P(t_0' - \tau_p) - P(t_0) + R(t_0' - \tau_r) - R(t_0) = 0. \quad (2.4)$$

By adding and subtracting $P(t_0') + R(t_0')$ and rearranging terms, (2.4) can also be written

$$[P(t_0') - P(t_0)] + [R(t_0') - R(t_0)] \\ = [P(t_0') - P(t_0' - \tau_p)] + [R(t_0') - R(t_0' - \tau_r)]. \quad (2.5)$$

For small values of τ_p and τ_r , such that $\delta_\tau = t_0' - t_0$ is sufficiently small, both sides of (2.8) can be represented in differential form as

$$\delta_\tau [P'(t_0) + R'(t_0)] = \tau_p P'(t_0) + \tau_r R'(t_0) \quad (2.6)$$

where $P'(t_0) = dP_0(t)/dt$ at $t = t_0$, and R' is correspondingly defined.

Equation (2.9) can be written in the form

$$\delta_\tau = p_\tau \tau_p + r_\tau \tau_r \quad (2.7)$$

where

$$p_\tau = \frac{P'(t_0)}{P'(t_0) + R'(t_0)}, \quad r_\tau = \frac{R'(t_0)}{P'(t_0) + R'(t_0)}, \quad (2.8)$$

and

$$p_\tau + r_\tau = 1. \quad (2.9)$$

With random uncorrelated displacements of rms values $\bar{\tau}_p$ and $\bar{\tau}_r$, the rms value of δ_τ is

$$\hat{\delta}_\tau = (p_\tau^2 \bar{\tau}_p^2 + r_\tau^2 \bar{\tau}_r^2)^{1/2} \quad (2.10)$$

Equation (2.9) and (2.10) give the timing deviations in regenerated pulses in terms of the deviations τ_p and τ_r in the received pulses and in the timing wave. To limit timing deviations in the regenerated pulses, it is necessary to make p_τ and the product $r_\tau \tau_r$ small. This will entail the use of a timing wave comparable in amplitude to that of the pulses, or greater, in conjunction with a small timing deviation τ_r in the timing wave.

2.3 Conversion of Amplitude Variations Into Timing Deviations

With partial retiming there is a conversion of amplitude variations in the received pulses and in the timing wave into timing deviations in the regenerated pulses.

Let the pulses have an amplitude variation a_p and the timing wave a_r expressed as fractions of the normal values. Pulses will then be regenerated at a time t_0' given by

$$(1 + a_p)P(t_0') + (1 + a_r)R(t_0') = L. \quad (2.11)$$

Subtracting (1.3) from (2.11),

$$[P(t_0') - P(t_0)] + [R(t_0') - R(t_0)] = -a_p P(t_0') - a_r P(t_0').$$

For small values of a_p and a_r , such that $\delta_a = t_0' - t_0$ is sufficiently small, the same procedure as in Section 2.2 gives

$$\delta_a = (p_a a_p + r_a a_r), \quad (2.12)$$

and

$$p_a = \frac{-P(t_0)}{P'(t_0) + R'(t_0)}, \quad r_a = \frac{-R(t_0)}{P'(t_0) + R'(t_0)}. \quad (2.13)$$

For uncorrelated variations of rms amplitude $\underline{a}_p$ and $\underline{a}_r$ the corresponding rms timing deviation is

$$\underline{\delta}_a = (p_a^2 \underline{a}_p^2 + r_a^2 \underline{a}_r^2)^{1/2}. \quad (2.14)$$

Equations (2.12) and (2.14) give the timing deviations in regenerated pulses resulting from amplitude variations in the pulses and in the timing wave.

2.4 Resultant Timing Deviations in Regenerated Pulses

For small variations in the pulses and in the timing wave as considered previously, the resultant timing deviation in a particular regenerated pulse is

$$\Delta = \delta_\tau + \delta_a. \quad (2.15)$$

Considering a large number of pulses, the resultant rms timing deviation in terms of the rms deviation in the received pulses and in timing wave is

$$\underline{\Delta} = (\underline{\delta}_\tau^2 + \underline{\delta}_a^2)^{1/2}. \quad (2.16)$$

These expressions can also be written

$$\Delta = \Delta_p + \Delta_r, \quad (2.17)$$

$$\underline{\Delta} = (\underline{\Delta}_p^2 + \underline{\Delta}_r^2)^{1/2}, \quad (2.18)$$

$$\Delta_p = p_\tau \tau_p + p_a a_p,$$

$$\underline{\Delta}_p^2 = p_\tau^2 \tau_p^2 + p_a^2 a_p^2, \quad (2.19)$$

$$\Delta_r = r_\tau \tau_r + r_a a_r,$$

$$\underline{\Delta}_r^2 = r_\tau^2 \tau_r^2 + r_a^2 a_r^2. \quad (2.20)$$

III ILLUSTRATIVE REGENERATION CHARACTERISTICS

3.0 General

In this section the general equations given in the preceding sections are applied to a particular case, in order to obtain specific expressions for the regeneration characteristics and illustrative curves, as an aid to further analysis. The particular case selected for illustration approximates the conditions in experimental Wrathall repeaters, and may be regarded as an idealized model of such a repeater, in which certain effects to be discussed later are ignored.

3.1 Pulse Shape

It will be assumed that the pulses are transmitted at intervals T and that the shape of the received pulses after equalization is given by:

$$P(t) = \frac{1}{2} \left[1 + \cos \frac{\pi t}{\eta T} \right]. \quad (3.1)$$

This is the familiar "raised cosine" type of pulse. With $\eta = 1$ the pulse width is the maximum that can be tolerated without intersymbol interference. With $\eta = \frac{3}{4}$, the amplitude of a pulse train at a point midway between two success pulses is equal to half the peak amplitude of a pulse. The latter assumption will be made here, for reasons discussed later.

3.2 Retiming Wave

The retiming wave is assumed to be given by

$$R(t) = -\frac{1}{2} \cos \psi \left[1 - \cos \left(2\pi \frac{t}{T} - \psi \right) \right]. \quad (3.2)$$

This type of retiming wave can be obtained if a sinusoidal wave of the pulse repetition frequency $f = 1/T$ is applied to a resonant circuit to reduce distortion of the timing wave by noise. The resonant circuit would have a nominal resonant frequency $f = 1/T$, but because of mistuning it would actually be f_0 . The output of the resonant circuit after appropriate adjustment of amplitude would be of the form [Appendix I, equation (2)]:

$$R_0(t) = \frac{1}{2} \cos \psi \cos \left(2\pi \frac{t}{T} - \psi \right), \quad (3.3)$$

where ψ is the phase shift of the resonant circuit at the frequency f ,

given by:

$$\tan \psi = Q \left(\frac{f}{f_0} - \frac{f_0}{f} \right), \quad (3.4)$$

and Q is the loss constant of the resonant circuit. If the peaks of the wave given by (3.3) are held at zero potential, a retiming wave as given by (3.2) is obtained. This type of retiming wave can also be obtained by applying an infinite sequence of rectangular pulses of equal amplitudes with spacing T to a resonant circuit.

3.3 Triggering Instants

With a pulse shape and retiming wave as assumed above, the resultant wave is given by

$$P(t) + R(t) = \frac{1}{2} \left[1 + \cos \frac{\pi t}{\eta T} \right] - \frac{\cos \psi}{2} \left[1 - \cos \left(2\pi \frac{t}{T} - \psi \right) \right]. \quad (3.5)$$

This wave is shown in Fig. 3 for $\psi = 0$ and $\pm 60^\circ$. For $\psi = \pm 90^\circ$ the retiming wave disappears, so that the combined wave is $P(t)$.

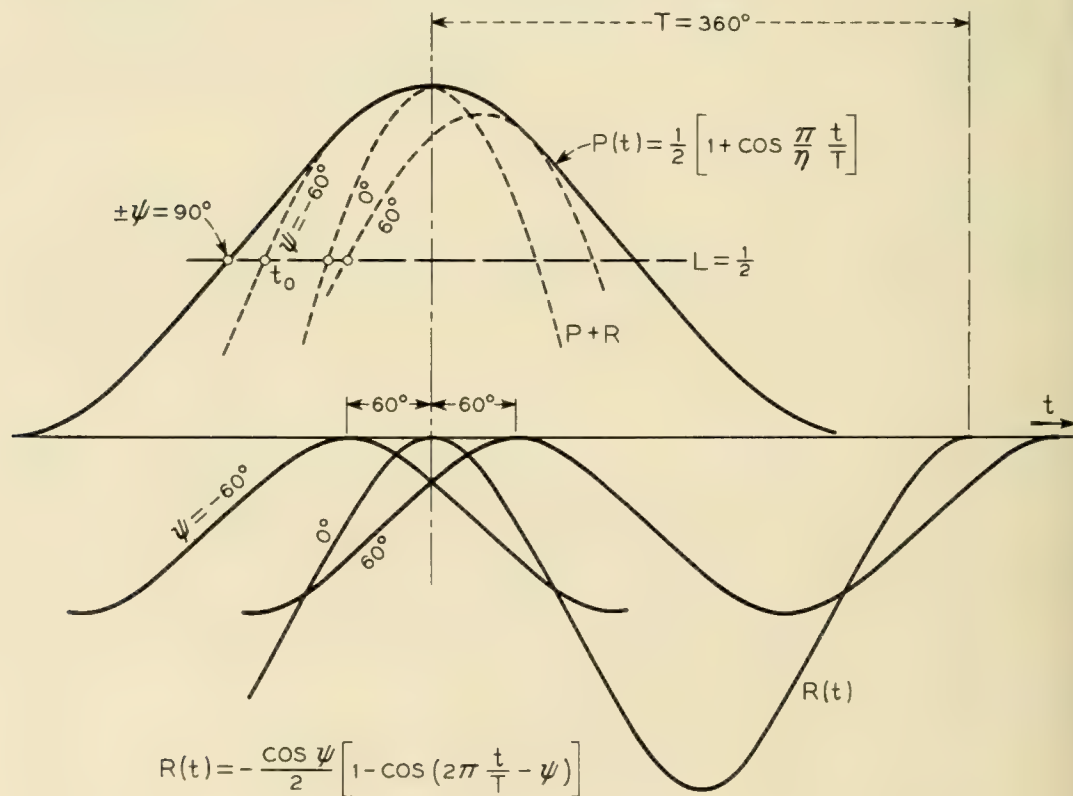


Fig. 3 — Illustrative example of pulse shape and retiming wave.

The triggering instants t_0 are obtained from the relation

$$P(t_0) + R(t_0) = L. \quad (3.6)$$

With complete retiming the optimum performance, with positive and negative noise amplitudes of equal probabilities, is obtained with a triggering level $\frac{1}{2}$. With partial retiming, optimum performance is obtained with a somewhat lower triggering level, but this is of secondary importance in connection with the present analysis. For this reason $L = \frac{1}{2}$ is assumed, in which case the following equation is obtained for determination of t_0 :

$$\cos \frac{\pi t_0}{\eta T} - \cos \psi \left[1 - \cos \left(2\pi \frac{t_0}{T} - \psi \right) \right] = 0. \quad (3.7)$$

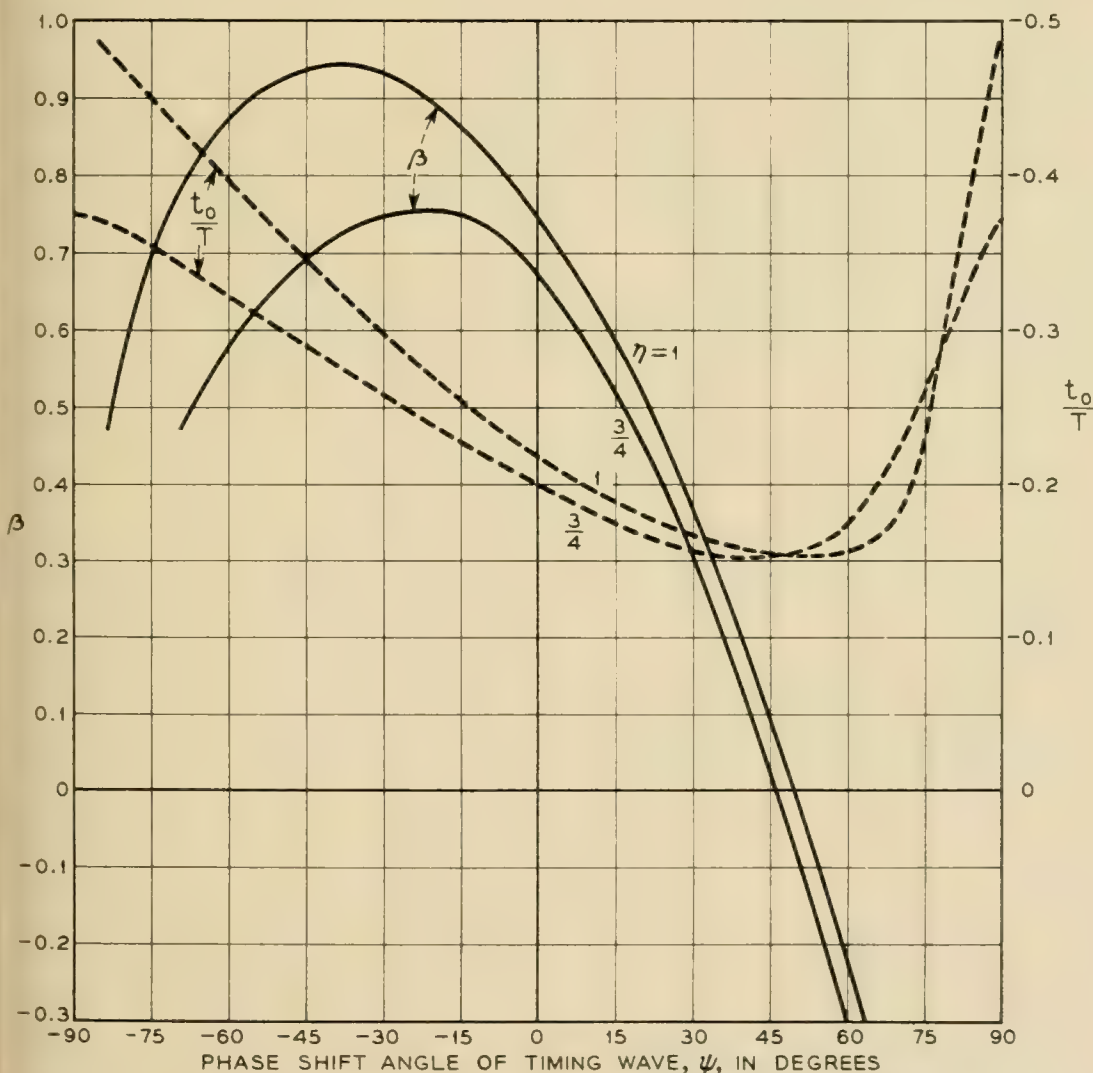


Fig. 4 — Triggering times versus phase shifts in timing wave.

TABLE I — VALUES OF t_0/T FOR $\eta = \frac{3}{4}$ AND $\eta = 1$

ψ	-90°	-60°	-30°	0	30°	60°	90°
$\eta = \frac{3}{4}$	-0.375	-0.322	-0.258	-0.198	-0.156	-0.17	-0.375
$\eta = 1$	-0.50	-0.391	-0.293	-0.215	-0.170	-0.15	-0.50

TABLE II — VALUES OF p_τ AND r_τ FOR $\eta = \frac{3}{4}$

ψ	-90°	-60°	-30°	0	30°	60°	90°
p_τ	1	0.61	0.43	0.32	0.32	0.50	1
r_τ	0	0.39	0.57	0.68	0.68	0.50	0

This equation is satisfied for the values of t_0/T given in Table I. The values of t_0/T are also shown in Fig. 4 as a function of ψ .

3.4 Conversion Factors for Time Deviations

The conversion factors defined by (2.8) become:

$$p_\tau = \frac{1}{D} \sin \frac{\pi}{\eta} \frac{t_0}{T} = 1 - r_\tau, \tag{3.8}$$

$$r_\tau = \frac{1}{D} 2\eta \cos \psi \sin \left(2\pi \frac{t_0}{T} - \psi \right), \tag{3.9}$$

and

$$D = \sin \frac{\pi}{\eta} \frac{t_0}{T} + 2\eta \cos \psi \sin \left(2\pi \frac{t_0}{T} - \psi \right), \tag{3.10}$$

where t_0/T has the values given previously as a function of ψ .

For various values of ψ , the factors for $\eta = \frac{3}{4}$ are given in Table II and in Fig. 5.

3.5 Conversion Factors for Amplitude Into Time Deviations

The conversion factors defined by (2.13) become

$$p_a = -\frac{T\eta}{\pi} \frac{1}{D} \left[1 + \cos \frac{\pi}{\eta} \frac{t_0}{T} \right], \tag{3.11}$$

and

$$r_a = \frac{T\eta}{\pi} \frac{1}{D} \cos \psi \left[1 - \cos \left(2\pi \frac{t_0}{T} - \psi \right) \right], \tag{3.12}$$

where D and t_0/T are defined as before.

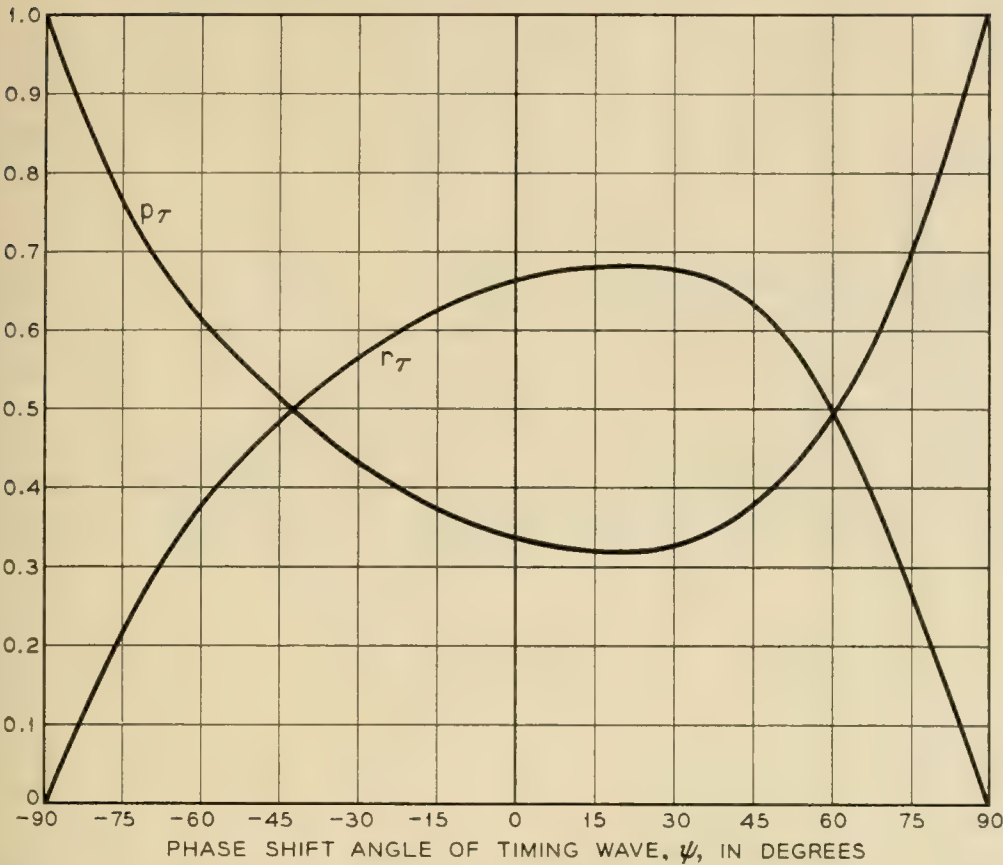


Fig. 5 — Conversion of timing deviations in received pulses and in timing wave into timing deviations in regenerated pulses, for pulse shapes and timing waves shown in Fig. 3. Timing deviations in regenerated pulses in relation to timing deviation t_p in pulses and t_r in retiming wave is $p_t t_p + r_t t_r$.

For various values of ψ the factors for $\eta = \frac{3}{4}$ are given in Table III and in Fig. 6.

3.6 Correlated Amplitude and Time Deviations

The amplitude and time deviations in the pulses are generally uncorrelated, but this does not always apply to the timing wave. In particular, if a deviation τ_r in the timing wave is the result of a change in the phase ψ , it will be accompanied by a given amplitude variation. A change in phase by $\Delta\psi$ is related to the corresponding time deviation τ_r by

$$\Delta\psi = \frac{2\pi}{T} \tau_r . \tag{3.13}$$

TABLE III — VALUES OF p_a/T AND r_a/T FOR $\eta = \frac{3}{4}$

ψ	-90°	-60°	-30°	0	30°	60°	90°
p_a/T	-0.24	-0.185	-0.175	-0.19	-0.22	-0.325	-0.24
r_a/T	0	0.035	0.055	0.072	0.106	0.14	0

With this change in phase, the factor $\cos \psi$ of (3.2) is modified to

$$\begin{aligned} \cos(\psi + \Delta\psi) &= \cos \psi \cos \Delta\psi - \sin \psi \sin \Delta\psi, \\ &\cong \cos \psi - \frac{2\pi}{T} \tau_r \sin \psi \end{aligned} \quad (3.14)$$

where the approximation applies for small values of ψ . The amplitude variation resulting from the above change in phase is accordingly

$$a_r = -\tau_r \frac{2\pi}{T} \sin \psi. \quad (3.15)$$

Considering both the time deviation τ_r and the corresponding amplitude variation a_r , the resultant time deviation in regenerated pulses is in accordance with (2.20)

$$\Delta_r = r_\tau \tau_r + r_a a_r. \quad (3.16)$$

The resultant equation can be written

$$\Delta_r = \beta \tau_r \quad (3.17)$$

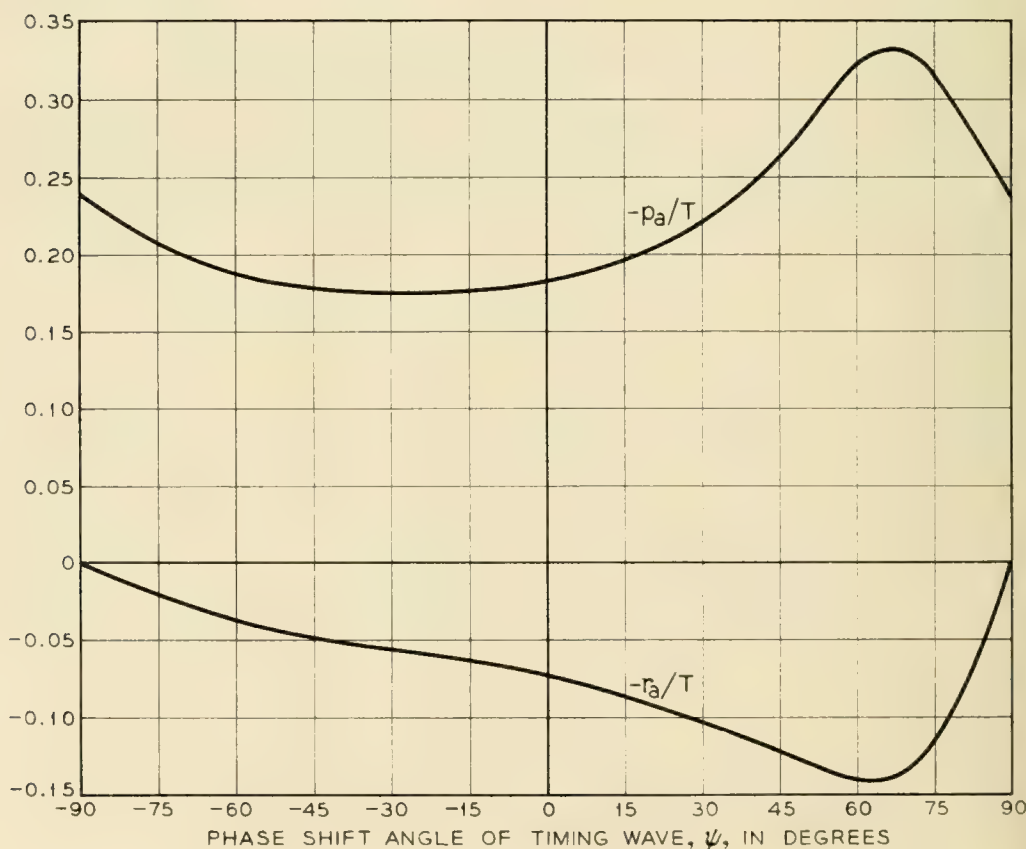


Fig. 6 — Conversion of amplitude variations in received pulses and in timing wave into timing deviations in regenerated pulses, for pulse shapes and timing waves shown in Fig. 3. Timing deviations in regenerated pulses for amplitude variations a_p and a_r in received pulses and in timing wave is $p_a a_p + r_a a_r$.

where

$$\beta = \frac{2\eta \cos \psi}{D} \left\{ \sin \left(2\pi \frac{t_0}{T} - \psi \right) + \sin \psi \left[1 - \cos \left(2\pi \frac{t_0}{T} - \psi \right) \right] \right\} \tag{3.18}$$

and D and t_0/T are defined as before.

The factor β indicates the time deviation in regenerated pulses in relation to the time deviation τ_r in the timing wave which results from a phase shift $\Delta\psi$ as given by (3.13). It may be regarded as a timing feedback factor that is of interest in connection with timing from regenerated pulses as discussed later. The factor β is shown in Fig. 4 for $\eta = \frac{3}{4}$ and $\eta = 1$.

3.7 Reduction in Tolerance to Noise by Timing Deviations

When the pulse shape is given by (3.1) and the timing wave is displaced by τ_0 , the tolerance to noise is in accordance with (2.2) reduced by the factor

$$\begin{aligned} \mu &= \frac{\frac{1}{2} \left(1 + \cos \frac{\pi \tau_0}{\eta T} \right) - \frac{1}{2}}{\frac{1}{2} \left(1 + \cos \frac{\pi 0}{\eta T} \right) - \frac{1}{2}} \\ &= \cos \frac{\pi \tau_0}{\eta T} . \end{aligned} \tag{3.19}$$

For a phase displacement ψ ,

$$\tau_0 = T\psi/2\pi, \tag{3.20}$$

and

$$\mu = \cos \frac{\psi}{2\eta}. \tag{3.21}$$

For $\eta = \frac{3}{4}$, the factor μ and the corresponding reduction in the tolerance to noise in db are as follows:

$\psi =$	0	$\pm 30^\circ$	$\pm 45^\circ$	$\pm 60^\circ$	$\pm 90^\circ$
$\mu =$	1	0.94	0.866	0.766	0.5
$\mu_{db} =$	0	0.5	1.2	2.3	6

IV DERIVATION OF TIMING WAVE FROM PULSE TRAIN

4.0 General

The retiming wave $R(t)$ must have a fixed relation to the received pulses, with certain tolerable fixed and random deviations to be considered later. Such a timing wave can be derived from the pulse train with the aid of a sufficiently narrow band-pass filter, the simplest form of which is a resonant circuit consisting of a coil and capacitor in series or in parallel.

A train of rectangular "on-off" pulses is shown in Fig. 7 as it would appear at the output of a regenerative repeater and at the input of the next repeater, (dotted) with uniform intervals T between sampling points.

As indicated in Fig. 7, the pulse train can be regarded as being made up of two components. One of these is an infinite sequence of pulses of one polarity, the other an infinite sequence of randomly positive and negative polarity.

It will be recognized that the first of the above components at the output has a fundamental frequency equal to the pulse repetition frequency, $f = 1/T$, and the forced response of a resonant circuit to this component will be the pulse repetition frequency, regardless of any imperfections in tuning. In order that this frequency be present in the received pulse train, it is necessary that the spectrum of the received pulses extend beyond the pulse repetition frequency, so that there will be a ripple in a long sequence of received pulses of one polarity, as indicated in the illustration.

The second random component of the pulse train will have a frequency spectrum that is nearly uniform over the band of the tuned circuit, and which will vary in amplitude depending on the composition of the pulse train. The response of the tuned circuit to this component is thus rather complex, and must be treated on an approximate statistical basis. It will consist of an almost periodic wave with random amplitude and phase modulation, and with mean frequency equal to the resonant frequency.

Owing to the presence of the second component, there will be a variation with time in the amplitude and phase of the response of a mistuned resonant circuit, and resultant deviations in timing. The regenerated pulses will thus not be uniformly spaced, but will in general have random deviations from the desired exact positions. Such deviations can be created by superposing on a train with uniform spacing a random dipulse train, as indicated in Fig. 7. The resonant circuit response to this

dipulse train would be expected to be smaller than to the random amplitude component of the pulse train. It may be regarded as a third component representing a second order effect resulting from the second component.

In the Appendix, this method of superposition has been used as a basis of an analysis of a resonant circuit response to a random binary pulse train. This problem has also been dealt with by somewhat different methods in prior unpublished work by W. R. Bennett and J. R. Pierce, both of Bell Telephone Laboratories.

In this analysis it is assumed that the regenerated pulses are of sufficiently short duration to be regarded as impulses. The response of the resonant circuit to the second and third components above, when taken in relation to that for the first component will, however, remain very

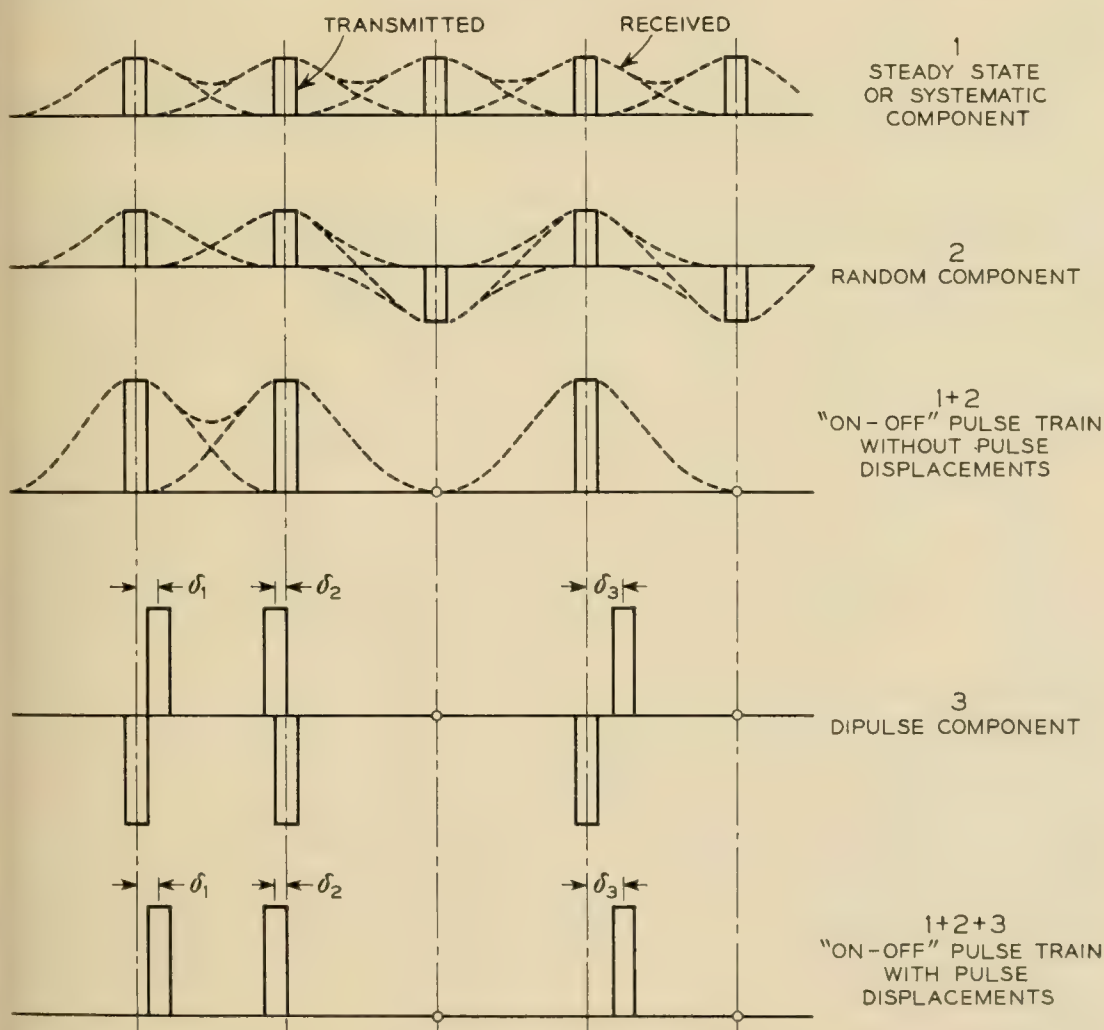


Fig. 7—Resolution of "on-off" pulse train with timing deviations into systematic component (1), random component (2), and time displacement component (3).

nearly the same for other pulse shapes, provided the frequency spectrum of the pulses can be regarded as approximately constant over the important portion of the band of the resonant circuit. This approximation is legitimate for resonant circuits with a loss constant Q and pulse shapes at the input of repeaters as considered here.

4.1 Resonant Circuit Response to Steady State Component

The first component consists of an infinite sequence of impulses of amplitude $\frac{1}{2}$ and all of the same polarity, at intervals T . This sequence has a fundamental frequency $f = 1/T$. When it impinges on a resonant circuit with resonant frequency $f_0 = f - \Delta f$ and loss constant Q , the response is of the form

$$A_s(t) = \cos \psi \cos (\omega t - \psi), \quad (4.1)$$

and

$$\tan \psi = Q(f/f_0 - f_0/f) \cong 2Q \frac{\Delta f}{f}. \quad (4.2)$$

The response is thus a steady state sinusoidal wave of frequency f displaced from the fundamental component of the input wave by the phase shift ψ and reduced in amplitude by $\cos \psi$. This is the phase shift and amplitude reduction of the resonant circuit at the frequency f when the resonant frequency is f_0 .

4.2 Resonant Circuit Response to Random Signal Component

The second component consists of an infinite random sequence of impulses of amplitude $\pm \frac{1}{2}$, at intervals T . The response of the resonant circuit to this component will be a randomly fluctuating wave $A_r(t)$ of mean value 0. The maximum positive amplitude is obtained when all impulses of the second component are positive and is $A_r(t) = A_s$. The maximum negative amplitude is $A_r(t) = -A_s$. Owing to the presence of this component the total output of the resonant circuit $A_s + A_r(t)$ can thus fluctuate between the limits 0 and $2A_s$, but the actual fluctuations of significant probability will be smaller.

The above fluctuations can be resolved into a component in phase with the steady state response given by (4.1) and another component at quadrature with the steady state timing wave. The rms values of these components taken in relation to the amplitude of the steady state wave are

$$\underline{a}_r' = \underline{A}_r'/A_s = \left(\frac{\pi}{2Q}\right)^{1/2} [1 - \psi^2/2]^{1/2} \frac{1}{\cos \psi}, \quad (4.3)$$

and

$$\underline{a}_r'' = \underline{A}_r''/A_s = \left(\frac{\pi}{4Q}\right)^{1/2} \frac{|\psi|}{\cos \psi}. \quad (4.4)$$

These relations apply for small values of ψ and for $\pi/Q \ll 1$.

The resultant rms amplitude variation in the timing wave is $\underline{a}_r = \underline{a}_r'$ as given by (4.3).

The rms phase error $\bar{\varphi}_r$ resulting from the quadrature component $\underline{a}_r''$ is given by

$$\tan \bar{\varphi}_r \cong \bar{\varphi}_r = \underline{a}_r''. \quad (4.5)$$

The corresponding rms time deviation is $(T/2\pi)\bar{\varphi}_r$ or

$$\underline{\delta}_r = \frac{T}{2\pi} \left(\frac{\pi}{4Q}\right)^{1/2} \frac{|\psi|}{\cos \psi}. \quad (4.6)$$

With regard to the probability of exceeding the above rms values by various factors the normal law can probably be invoked with reasonable accuracy. As mentioned before, the maximum possible amplitudes are $\hat{A}_r(t) = \pm A_s$ which would correspond to a peak factor $(2Q/\pi)^{1/2}$. With $Q = 100$, the factor is about 8, while with $Q = 1000$ it is about 25. Based on the normal law the probability of exceeding the rms value by a factor of 4 is about 5×10^{-5} , and by a factor of 5, about 10^{-7} . The normal law would be expected to apply, since the limiting peak values are substantially greater than the peak values expected with significant probabilities.

4.3 Resonant Circuit Response to Pulse Displacements

Because of the random components given by (4.3) and (4.4), the timing wave will contain small random amplitude and phase deviations from a sinusoidal wave represented by (4.1). This will result in small random deviations in the positions of regenerated pulses triggered from the timing wave, which is represented by the third component shown in Fig. 7. When the rms deviation in the pulse positions is $\underline{\delta}$, there will be an additional random quadrature component in the timing wave which, when taken in relation to the steady state component, is given by

$$\underline{a}_\delta'' = \underline{A}_\delta''/A_s = \omega \underline{\delta} \left(\frac{\pi}{Q}\right)^{1/2}. \quad (4.7)$$

The corresponding rms phase deviation is given by

$$\bar{\varphi}_\delta \cong \underline{a}_\delta''. \quad (4.8)$$

The resultant rms time deviation is $(T/2\pi)\bar{\varphi}_\delta$ or

$$\delta_\delta = \delta\alpha, \quad (4.9)$$

and

$$\alpha = (\pi/Q)^{1/2}. \quad (4.10)$$

The above factor α applies to a single resonant circuit. When the rms timing deviations represented by (4.9) are present in the regenerated pulse train, the rms deviation at the output of the second resonant circuit is

$$\delta_{\delta,2} = \delta\alpha_1\alpha_2,$$

where

$$\alpha_1 = \alpha.$$

With n resonant circuits in tandem,

$$\delta_{\delta,n} = \delta\alpha_1\alpha_2\alpha_3 \cdots \alpha_n. \quad (4.11)$$

The factors α_n are given by

$$\alpha_1 = \alpha = (\pi/Q)^{1/2}, \quad (4.12)$$

$$\alpha_j = \left(1 - \frac{1}{2(j-1)}\right)^{1/2} \quad j \geq 2, \quad (4.13)$$

$$\alpha_2 = (1 - \frac{1}{2})^{1/2},$$

$$\alpha_3 = (1 - \frac{1}{4})^{1/2},$$

$$\alpha_4 = (1 - \frac{1}{6})^{1/2}, \text{ etc.}$$

$$[\alpha_1 \cdot \alpha_2 \cdot \alpha_3 \cdots \alpha_n]^2 = \alpha^2 \frac{1 \cdot 3 \cdot 5 \cdots [2(n-1) - 1]}{2 \cdot 4 \cdot 6 \cdots 2(n-1)} \quad (4.14)$$

$$= \alpha^2 \frac{(2n)!}{2^{2n}(n!)^2} \quad (4.15)$$

$$= \alpha^2 \left(\frac{1}{\pi n}\right)^{1/2} \quad \text{when} \quad n \gg 1. \quad (4.16)$$

The factors α_j for $j \geq 2$ represent the reduction in timing deviations resulting from the reduction in bandwidth as resonant circuits are added in tandem. If resonant circuits with a narrow flat pass-band were used, the bandwidth of any number of resonant circuits in tandem would be the same as for a single resonant circuit. In this case $\alpha_2 = \alpha_3 = \alpha_n = 1$.

4.4 Deviations in Timing Wave

The timing wave derived from an "on-off" pulse train with the aid of a resonant circuit will in accordance with the expressions given in the previous sections contain three types of amplitude and timing deviations.

The first type is a fixed amplitude reduction by a factor a_0 and a fixed time deviation τ_0 given by

$$a_0 = \cos \psi, \quad (4.17)$$

and

$$\tau_0 = \frac{T}{2\pi} \psi, \quad (4.18)$$

where ψ is given by (4.2).

The second type is a random amplitude and time deviation resulting from the random amplitude component of the pulse train, which have rms values

$$a_r \cong \left(\frac{\pi}{2Q} \right)^{1/2} [1 - \psi^2/2]^{1/2} \frac{1}{\cos \psi}, \quad (4.19)$$

and

$$\delta_r \cong \frac{T}{2\pi} \left(\frac{\pi}{4Q} \right)^{1/2} \frac{|\psi|}{\cos \psi}. \quad (4.20)$$

The third type is a random amplitude and time deviation resulting from random timing deviations $\bar{\tau}_p = \delta$ in the pulse train. The amplitude variation can be disregarded and the rms time deviation is

$$\delta_\delta = \alpha \bar{\tau}_p, \quad \alpha = \left(\frac{\pi}{Q} \right)^{1/2}. \quad (4.21)$$

The total rms amplitude variation is accordingly given by (4.19). The total rms timing deviation obtained by combining (4.20) and (4.21) is

$$\bar{\tau}_r = (\delta_r^2 + \alpha^2 \bar{\tau}_p^2)^{1/2}. \quad (4.22)$$

The expressions for δ_r and $\bar{\tau}_r$ are the quantities appearing in (2.20) for Δ_r , the total rms timing deviation in regenerated pulses resulting from random amplitude and timing deviations in the timing wave.

V SELF-TIMED REPEATERS WITH PARTIAL RETIMING

5.0 General

As shown in the preceding section, timing for pulse regeneration can be derived from the pulse trains, with certain random phase and amplitude variations in the timing wave that can be reduced by increasing the loss constant Q of the resonant circuit. This method of "self-timing" can be combined with partial retiming, and the regeneration characteristics of this type of repeater will be discussed in the following sections.

For purposes of numerical illustration, the same type of pulse shape and timing wave will be assumed as in the previous numerical illustration in Section III. This pulse shape and timing wave closely approximates those in experimental Wrathall repeaters, in which timing is derived from the regenerated pulse train. In the following discussion timing from the received pulse train will also be considered.

5.1 Timing from Received Pulse Train

It will be assumed that the timing wave is derived from the received pulse train with the aid of a resonant circuit and that random timing deviations are absent. The response of the resonant circuit is then a sinusoidal wave as given by (4.1). From this wave it is possible to obtain a retiming wave of the form

$$R(t) = -\cos \psi \left[1 - \cos \left(2\pi \frac{t}{T} - \psi \right) \right]. \quad (5.1)$$

This can be accomplished by holding the peaks of the timing wave from the resonant circuit at zero potential with a diode. This is the form of retiming wave previously considered in Section III, in conjunction with a pulse shape given by (3.1).

As shown in Section 3.7, the tolerance to noise will vary with the phase shift ψ of the resonant circuit, in accordance with (3.21). If a reduction in the tolerance to noise of about 2 db is allowed, the maximum permissible phase shift would be about $\psi = 1$ radian (57.6°). On this basis the maximum permissible deviation $\Delta f_{\max}$ in the resonant frequency from the pulse repetition frequency f as obtained from (4.2) with $\psi = 1$ radian becomes

$$\frac{\Delta f_{\max}}{f} = \frac{\tan \psi}{2Q} = \frac{1.58}{2Q}. \quad (5.2)$$

For various values of Q in the range that can be realized by simple

resonant circuits, the permissible deviations are as follows:

Q	10	25	50	100	200
$\Delta f_{\max}/f$	0.08	0.030	0.016	0.008	0.004

This assumes that there are no random timing deviations and that the tolerance to noise is reduced by not more than 2 db.

5.2 Timing from Regenerated Pulse Train

It will again be assumed that there are no random timing deviations. Without a phase shift in the resonant circuit, let the regenerated pulses be triggered at a time t_0 . When there is a phase shift ψ' , the pulses will be triggered at a time t_0' . The timing wave derived from the regenerated pulses will then have a time shift

$$\Delta = t_0' - t_0 + \frac{T}{2\pi} \psi'.$$

This time shift will cause pulses to be regenerated with a time shift $\beta' \Delta$, which must equal $t_0' - t_0$. Accordingly,

$$t_0' - t_0 = \beta' \left(t_0' - t_0 + \frac{T}{2\pi} \psi' \right),$$

and

$$t_0' - t_0 = \frac{T}{2\pi} \frac{\beta' \psi'}{1 - \beta'}. \quad (5.3)$$

With timing from the received pulse train with a phase shift ψ in the resonant circuit, the following relation applies:

$$t_0' - t_0 = \frac{T}{2\pi} \beta \psi. \quad (5.4)$$

If $t_0' - t_0$ is to be the same in both cases, so that the timing wave and tolerance to noise is the same, the following relation must exist between the phase shifts in the resonant circuit:

$$\psi' = \psi (1 - \beta') \frac{\beta}{\beta'}. \quad (5.5)$$

In this expression, β and β' are the factors shown in Fig. 4. It will be recognized from (5.5) that the smallest permissible phase shifts are obtained for large values of β' . From Fig. 4, it is seen that the largest

values of β are for phase shifts between 0 and -60° . For $\eta = \frac{3}{4}$, $\beta \cong 0.7$ and for $\eta = 1$, $\beta \cong 0.9$.

For $\eta = \frac{3}{4}$ and $\eta = 1$ the tolerable maximum phase shifts ψ' in the resonant circuit with timing from the regenerated pulse train, in relation to the maximum tolerable ψ with timing from the input, are

$$\psi' \cong 0.3\psi \quad \text{for} \quad \eta = \frac{3}{4}, \quad (5.6)$$

and

$$\psi' \cong 0.1\psi \quad \text{for} \quad \eta = 1.$$

Although greater phase shifts can be tolerated when ψ is positive, and β' is smaller than above, the requirements on the resonant circuit must be based on the worst condition that can be encountered, as above.

From (5.6) it follows that for $\eta = 1$ the requirements on the permissible phase shift in the resonant circuit are much more severe than for $\eta = \frac{3}{4}$. For this reason the latter value of η is decidedly preferable for the particular case in which the peak amplitudes of the pulse train and the timing waves are equal, as assumed here. A value $\eta = \frac{3}{4}$ is also desirable from the standpoints of avoiding intersymbol interference between adjacent pulses at the triggering instants, to permit the timing wave to be derived from the pulse train and to permit self-starting of the repeaters, as discussed later.

In accordance with (5.6) the maximum tolerable frequency deviation for $\eta = \frac{3}{4}$ will be less than with timing from the received pulse train by a factor of about 0.3. The maximum permissible frequency deviation for a phase shift of about one radian in the timing wave and 0.3 radian in the resonant circuit, will accordingly be about as follows:

Q	10	25	50	100	200
$\Delta f_{\max}/f$	0.025	0.009	0.005	0.0025	0.0012

For a repeater with complete rather than partial retiming, the factor β would be unity, and timing from the regenerated pulse train would not be possible.

5.3 Random Timing Deviations

In combining random timing deviations from various sources at a particular repeater, it will be assumed that there is no correlation between the various deviations, so that they will combine on a root-sum-square basis.

In accordance with (2.21) the rms timing deviation at the output is then:

$$\underline{\Delta}^2 = (p_\tau^2 \bar{\tau}_p^2 + p_a^2 \underline{a}_p^2) + (r_\tau^2 \bar{\tau}_r^2 + r_a^2 \underline{a}_r^2), \quad (5.7)$$

where in accordance with (4.13) and (4.16)

$$\underline{a}_r = \left[\frac{\pi}{2Q} (1 - \psi^2/2) \right]^{1/2} \frac{1}{\cos \psi}, \quad (5.8)$$

$$\bar{\tau}_r = (\hat{\delta}_r^2 + \alpha^2 \bar{\tau}_p^2)^{1/2}, \quad (5.9)$$

$$\alpha = \left(\frac{\pi}{Q} \right)^{1/2}, \quad (5.10)$$

$$\hat{\delta}_r = \frac{T}{2\pi} \left(\frac{\pi}{4Q} \right)^{1/2} \frac{\psi}{\cos \psi}. \quad (5.11)$$

When (5.9) is inserted in (5.7)

$$\underline{\Delta}^2 = (p_\tau^2 + \alpha^2 r_\tau^2) \bar{\tau}_p^2 + p_a^2 \underline{a}_p^2 + r_\tau^2 \hat{\delta}_r^2 + r_a^2 \underline{a}_r^2. \quad (5.12)$$

This expression gives the rms timing deviation at the output in terms of the rms deviation $\bar{\tau}_p$ at the input and the various repeater parameters.

With timing from the output, rather than the input as assumed above, $\bar{\tau}_p$ is replaced by $\underline{\Delta}$ in (5.9), and the following relation is obtained:

$$\underline{\Delta}^2 (1 - \alpha^2 r_\tau^2) = p_\tau^2 \bar{\tau}_p^2 + p_a^2 \underline{a}_p^2 + r_\tau^2 \hat{\delta}_r^2 + r_a^2 \underline{a}_r^2. \quad (5.13)$$

In the above expressions $p_\tau^2 \cong 0.15$, $r_\tau^2 \cong 0.4$ and $\alpha^2 \cong 0.03$ ($Q = 100$). The term $\alpha^2 r_\tau^2$ can thus be neglected in comparison with p_τ^2 in (5.12) and in comparison with 1 in (5.13).

The following expression is thus obtained with timing from either the input or the output:

$$\begin{aligned} \underline{\Delta}^2 &= (p_\tau^2 \bar{\tau}_p^2 + p_a^2 \underline{a}_p^2) + (r_\tau^2 \hat{\delta}_r^2 + r_a^2 \underline{a}_r^2) \\ &= \underline{\Delta}_p^2 + \underline{\Delta}_r^2. \end{aligned} \quad (5.14)$$

5.4 Magnitude of Random Timing Deviations

The first two terms of (5.14) represents the rms timing deviations in the regenerated pulses resulting from timing deviations and amplitude variations in the received pulses. The last two terms represent the timing deviations resulting from timing deviations and amplitude variations in the timing wave. The conversion factors p_τ , p_a , r_τ and r_a are discussed in Section II and representative values given in Figs. 5 and 6. The values of $\underline{a}_r$ and $\hat{\delta}_r$ are obtained from (5.8).

TABLE IV — RMS DEVIATIONS FROM TIMING WAVE
DISTORTION FOR $Q = 100$

ψ	-60°	-30°	0°	30°	60°
$r_r \underline{\delta}_r / T$	0.011	0.005	0	0.006	0.015
$r_a \underline{a}_r / T$	0.006	0.007	0.009	0.014	0.024
$\underline{\Delta}_r / T$	0.0126	0.009	0.009	0.015	0.028
$\bar{\varphi}_r$	4.5°	3.2°	3.2°	5.4°	10°

In Table IV are given the values of the two last terms in (5.14), which represents the rms deviations $\underline{\Delta}_r$ owing to random deviations in the timing wave. The results are given for the particular case in which $Q = 100$, and for other values of Q are inversely proportional to $Q^{1/2}$. The table shows the deviations as a fraction of the interval T between pulses, and also as the corresponding rms phase deviation $\bar{\varphi}_r$.

In Table V are given the values of the first two terms in (5.14), which represents the rms deviation $\underline{\Delta}_p$ in the regenerated pulses resulting from random amplitude and timing deviation in the received pulses. In binary systems it is customary to limit the rms pulse distortion to $\underline{a}_p = \frac{1}{10}$, corresponding to $\frac{1}{10}$ the peak amplitude of the received pulses, or $\frac{1}{5}$ the triggering level (17 db signal-to-noise ratio). The corresponding rms phase deviation would be about $\frac{1}{10}$ radian, corresponding to an rms deviation $\bar{\tau}_p$ in the pulses of 0.016 the pulse spacing, or $\bar{\tau}_p / T \cong 0.016$. The total rms timing deviation obtained from (5.14) and the corresponding rms phase deviation are given in Table VI.

TABLE V — RMS DEVIATIONS RESULTING FROM PULSE DISTORTION

ψ	-60°	-30°	0°	30°	60°
$p_a \underline{a}_p / T$	0.019	0.018	0.019	0.022	0.032
$p_r \bar{\tau}_p / T$	0.010	0.007	0.005	0.005	0.008
$\underline{\Delta}_p / T$	0.021	0.020	0.020	0.023	0.033
$\bar{\varphi}_p$	7.5°	7.2°	7.2°	8.2°	12°

TABLE VI — TOTAL RMS DEVIATIONS FROM TIMING WAVE
AND PULSE DISTORTION

ψ	-60°	-30°	0°	30°	60°
$\underline{\Delta} / T$	0.025	0.022	0.022	0.028	0.043
$\bar{\varphi}$	9°	8°	8°	10°	16°

The probability that random phase deviations will exceed the above rms values by a factor of more than 4 is small enough to be ignored. On this basis the sum of the fixed and random deviations would be limited to about 70° , if the fixed phase shift ψ is less than $\pm 30^\circ$. With this requirement on the fixed phase shift for satisfactory performance, the values of $\Delta f_{\max}$ would be about half as great as previously given in Sections 5.1 and 5.2, for a single repeater as considered here.

VI REPEATER CHAINS

6.0 General

In the previous section, a single self-timed repeater was considered, from the standpoint of fixed and random timing deviations, as determined by various repeater design parameters. In a repeater chain there will be some cumulation of random timing deviations as the number of repeaters in tandem is increased, and a resultant reduction in the tolerance to noise of repeaters toward the end of the chain. Exact evaluation of such cumulation is rendered difficult by the circumstance that timing deviations from various sources may not follow the same law of combination along the repeater chain. In the following, expressions are given based both on root-sum-square and direct addition of random timing deviations, which can be regarded as lower and upper limits.

6.1 Combination of Random Timing Deviations

To determine the rms value of random timing deviations at the end of a repeater chain, it is necessary to combine random deviations from various repeaters. Random deviations from various sources at a repeater do not necessarily follow the same law of cumulation along a repeater chain. Since there is no correlation between timing deviations caused by noise in various repeater sections, these can be combined on a root-sum-square basis. This, however, may not be appropriate as regards the combination of timing deviations resulting from imperfections in the timing wave. Thus, with perfect tuning of all resonant circuits, the timing waves at various repeaters would have virtually identical amplitude variations, but no phase deviations. While in this case there would be complete correlation between the timing wave variations at the repeaters, it does not follow that the resultant timing deviations should be combined directly rather than on a root-sum-square basis along the repeater chain. The timing deviations at the end of a chain of N repeaters resulting from amplitude variations in the timing wave of the first repeater will be modified by N intermediate resonant circuits. Those

resulting from amplitude variations at subsequent repeaters will be modified by $N-1$, $N-2$ etc. intermediate resonant circuits. The situation is similar to that of applying identical noise waves at the input of each of N resonant circuits in tandem. At the output the N noise waves will have different shapes owing to restriction of the band and increasing phase distortion as the number of resonant circuits in tandem increases. For this reason combination on a root-sum-square basis appears justified also in this case, particularly with various degrees of mistuning of the resonant circuits, so that the amplitude variations in the timing waves will differ in phase among repeaters.

6.2 Propagation of Timing Deviations

To determine the cumulation of timing deviations along a repeater chain, it is convenient to first consider a single repeater as a source of timing deviations, and to determine the propagation of these timing deviations along a repeater chain. In the following, γ_n will designate the rms propagation factor for n repeaters in tandem; i.e., the factor by which the rms timing deviations at the end of a chain of n repeaters is smaller than at the first repeater, with timing deviations originating at the first repeater only.

Let the rms timing deviation at the output of the first repeater as given by (5.14) for convenience be taken as unity. At the output of the second repeater the squared rms timing deviation is then reduced by the factor

$$\gamma_2^2 = p_\tau^2 + \alpha_1^2 r_\tau^2, \quad \alpha_1 = \alpha. \quad (6.1)$$

As indicated symbolically in Fig. 8, the first term represents the reduction owing to partial retiming. The second term is the additional deviation

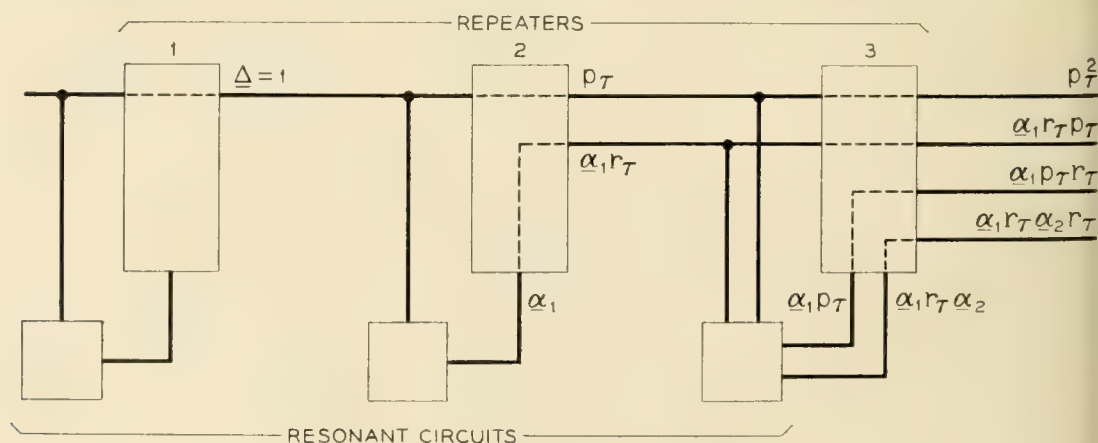


Fig. 8 — Propagation of random timing deviations along repeater chain.

tion resulting from the effect on the timing wave of unit rms deviation in the received pulse train at the second repeater.

At the output of the third repeater, the squared rms deviation is smaller than at the output of the first repeater by the factor

$$\gamma_3^2 = (p_\tau^2 + \alpha_1^2 r_\tau^2) p_\tau^2 + p_\tau^2 \alpha_1^2 r_\tau^2 + \alpha_1^2 r_\tau^2 \alpha_2^2 r_\tau^2 \quad (6.2)$$

$$= p_\tau^4 + 2\alpha_1^2 p_\tau^2 r_\tau^2 + \alpha_1^2 \alpha_2^2 r_\tau^4. \quad (6.3)$$

As indicated in Fig. 8, the first term in (6.2) represents the reduction owing to partial retiming of the received pulse train at the third repeater. The second term, $(p_\tau \alpha_1 r_\tau)^2$, is the additional rms deviation resulting from the effect on the timing wave at the third repeater of an rms deviation p_τ in the received pulse train. The third term $(\alpha_1 r_\tau \cdot \alpha_2 r_\tau)^2$ is the additional deviation caused by the effect on the timing wave of an rms deviation $\alpha_1 r_\tau$ in the received pulse train. The factor $\alpha_2 r_\tau$ represents the modification in the rms deviation $\alpha_1 r_\tau$ by a second resonant circuit, with α_2 defined as in Section 4.3.

At the output of the fourth repeater, the rms timing deviation is reduced by the following factor, obtained in the same manner:

$$\gamma_4^2 = p_\tau^6 + 3\alpha_1^2 p_\tau^4 r_\tau^2 + 3\alpha_1^2 \alpha_2^2 p_\tau^2 r_\tau^4 + \alpha_1^2 \alpha_2^2 \alpha_3^2 r_\tau^6. \quad (6.4)$$

At the output of repeater n , the squared rms timing deviation is smaller than at the output of the first repeater by the propagation factor

$$\begin{aligned} \gamma_n^2 = & p_\tau^{2(n-1)} + \frac{(n-1)}{1!} p_\tau^{2(n-2)} r_\tau^2 \alpha_1^2 \\ & + \frac{(n-1)(n-2)}{2!} p_\tau^{2(n-3)} r_\tau^4 \alpha_1^2 \alpha_2^2 \\ & + \frac{(n-1)(n-2)(n-3)}{3!} p_\tau^{2(n-4)} r_\tau^6 \alpha_1^2 \alpha_2^2 \alpha_3^2 \\ & + \cdots + r_\tau^{2(n-1)} \alpha_1^2 \alpha_2^2 \alpha_3^2 \cdots \alpha_{n-1}^2, \end{aligned} \quad (6.5)$$

where p_τ and r_τ are defined as in Section 2.2, and $\alpha_1, \alpha_2 \cdots \alpha_n$ as in Section 4.3.

In the above formulation the rms deviation at the output of the first repeater was assumed given by (5.14), which is an approximation of (5.12) in which the term $\alpha^2 r_\tau^2 \bar{\tau}_p^2$ was neglected. This term will have a different propagation factor ρ_n , the expression for which differs from that for γ_n as given by (6.5) in that the subscripts of the factors α_j are raised by one unit. Thus,

$$\rho_n^2 = p_{\tau}^{2(n-1)} + \frac{(n-1)}{1!} p_{\tau}^{2(n-2)} r_{\tau}^2 \alpha_2^2 + \cdots + r_{\tau}^{2(n-1)} \alpha_2^2 \alpha_3^2 \cdots \alpha_n^2. \quad (6.6)$$

The rms deviation at the output of repeater n thus becomes

$$\Delta_n^2 = (\Delta_p^2 + \Delta_r^2) \gamma_n^2 + \alpha^2 r_{\tau}^2 \bar{\tau}_p^2 \rho_{\tau}^2. \quad (6.7)$$

In the case of repeaters with partial retiming the last term in (6.7) can be neglected, in which case the cumulation of timing deviation will be virtually the same when the timing wave is derived from the regenerated as when it is derived from the received pulse train.

The above expressions apply for resonant circuits consisting of a coil and capacitor which have a gradual cut-off. If resonant circuits with a flat pass-band and sharp cut-offs were used, $\alpha_2 = \alpha_3 = \alpha_n$ and (6.5) can be simplified to

$$\gamma_n^2 = (1 - \alpha_1^2) p_{\tau}^{2(n-1)} + \alpha_1^2 (p_{\tau}^2 + r_{\tau}^2)^{(n-1)}. \quad (6.8)$$

6.3 Cumulation of Timing Deviations

The cumulation of random timing deviations from various repeaters in a chain can be determined from the propagation constant given above for any prescribed law of combination of timing deviations from various repeaters. When equal rms deviations are contributed by each of N repeaters, and they are combined on a root-sum-square basis, the rms deviation at the end of a repeater chain is greater than for a single repeater by the cumulation factor

$$C = \left(\sum_{n=1}^N \gamma_n^2 \right)^{1/2}. \quad (6.9)$$

An upper limit to C is obtained by taking $\alpha_2 = \alpha_3 = \alpha_n = 1$ in (6.5) in which case γ_n^2 is given by (6.8); (6.9) then becomes for $N = \infty$

$$C = \left[(1 - \alpha_1^2) \frac{1}{1 - p_{\tau}^2} + \alpha_1^2 \frac{1}{1 - p_{\tau}^2 - r_{\tau}^2} \right]^{1/2} \quad (6.10)$$

$$\cong \left(\frac{1}{1 - p_{\tau}^2} \right)^{1/2}, \quad (6.11)$$

where the terms in α_1^2 have been neglected in (6.11), since $\alpha_1^2 = \alpha^2 \ll 1$, about 0.03 for $Q = 100$.

From Fig. 5 it will be seen that when $\psi < \pm 60^\circ$, $p_{\tau} < 0.6$. Hence $C < 1.25$. Cumulation of random timing deviations can thus for practical

purposes be disregarded, with root-sum-square combination as assumed above. The value of C obtained from (6.11) will differ from that obtained from (6.9) when γ_n is given by (6.5), by a small fraction of one per cent.

Although root-sum-square combination appears justified for reasons given before, it is of interest to determine an upper limit to the cumulation based on direct addition of random timing deviations. The maximum cumulation factor thus obtained is

$$C_{\max} = \sum_{n=1}^N \gamma_n. \quad (6.12)$$

Employing (6.8) for γ_n and neglecting the terms in α_1^2 , the upper limit to the cumulation factor for $N = \infty$ becomes

$$C_{\max} = \frac{1}{1 - p_r}. \quad (6.13)$$

With $p_r < 0.6$ for $\psi < \pm 60^\circ$, $C_{\max} < 2.5$.

If the above maximum cumulation factor is applied to random timing deviations resulting from amplitude variations in the timing wave, as given in Table IV of Section 5.4, the resultant rms phase deviation at the end of a long repeater chain could be as great as 25° , rather than 10° for a single repeater, when $\psi = 60^\circ$ and $Q = 100$. To attain satisfactory performance it would in this case be necessary to limit the maximum fixed phase shift to substantially less than $\pm 60^\circ$, which would entail greater frequency precision than indicated in Sections 5.1 and 5.2.

If $\psi < \pm 15^\circ$, $p_r < 0.40$ and $C_{\max} < 1.7$. In this case the rms phase deviation as given in Table I for a single repeater is $\bar{\varphi}_r \cong 4^\circ$, and the rms phase deviation in a long repeater chain would be less than 7° . In a long repeater chain the rms phase deviation resulting from pulse distortion would be greater than given in Table II by an rms cumulation factor $C = 1.08$ for $p_r = 0.4$, and would thus be about 8° when $\psi < \pm 15^\circ$. The total rms phase deviation would thus be about $(7^2 + 8^2)^{1/2} \cong 11^\circ$. Random phase deviations exceeding 4 times the latter value, or about 45° , would be rather unlikely. The sum of the fixed and random phase deviations would thus be limited to about 60° , so that satisfactory performance would be expected when the fixed phase deviation is limited to about $\pm 15^\circ$.

With the approximations for γ_n employed above, the rms cumulation factor for a chain of N repeaters as obtained from (6.9) is less than for $N = \infty$ by the factor $(1 - p_r^{2N})^{1/2} \cong 0.99$ for $p_r = 0.5$ and $N = 3$. The maximum cumulation factor obtained from (6.12) is less than for $N = \infty$ by the factor $1 - p_r^N \cong 0.99$ for $N = 6$. Thus, cumulation of random

timing deviations is virtually completed in a chain of 3 to 6 repeaters, so that for experimental determinations of the degree of cumulation it suffices to operate a few repeaters in tandem.

6.4 Repeaters with Complete Retiming

In the particular case of complete retiming, $p_r = 0$ and $r_r = 1$ in (6.5) and (6.6) so that

$$\gamma_n = \alpha_1 \alpha_2 \alpha_3 \cdots \alpha_{n-1}, \quad (6.14)$$

$$\rho_n = \alpha_2 \alpha_3 \cdots \alpha_n. \quad (6.15)$$

For $n \gg 1$, approximation (4.16) can be employed, in which case

$$\gamma_n = \alpha \left(\frac{1}{\pi n} \right)^{1/4}, \quad \rho_n = \left(\frac{1}{\pi n} \right)^{1/4}. \quad (6.16)$$

In this case (5.14) simplifies to

$$\underline{\Delta}_p^2 + \underline{\Delta}_r^2 = \underline{\delta}_r^2, \quad (6.17)$$

since $p_a = 0$, $r_a = 0$, $p_r = 0$ and $r_r = 1$.

Hence (6.7) becomes

$$\underline{\Delta}_n^2 = \underline{\delta}_r^2 \gamma_n^2 + \bar{\tau}_p^2 \alpha^2 \rho_n^2. \quad (6.18)$$

With approximations (6.16),

$$\underline{\Delta}_n^2 = (\underline{\delta}_r^2 + \bar{\tau}_p^2) \alpha^2 \left(\frac{1}{\pi n} \right)^{1/2}. \quad (6.19)$$

At the output of the first repeater,

$$\underline{\Delta}_1^2 = \underline{\delta}_r^2 + \alpha^2 \bar{\tau}_p^2. \quad (6.20)$$

For $n \gg 1$ the squared propagation factor is accordingly

$$\underline{\Delta}_n^2 / \underline{\Delta}_1^2 = \alpha^2 \frac{\underline{\delta}_r^2 + \bar{\tau}_p^2}{\underline{\delta}_r^2 + \alpha^2 \bar{\tau}_p^2} \left(\frac{1}{\pi n} \right)^{1/2}. \quad (6.21)$$

The squared rms cumulation factor for $N \gg 2$ repeaters becomes

$$\begin{aligned} C^2 &\cong 1 + \alpha^2 \frac{\underline{\delta}_r^2 + \bar{\tau}_p^2}{\underline{\delta}_r^2 + \alpha^2 \bar{\tau}_p^2} \int_2^N \left(\frac{1}{\pi n} \right)^{1/2} dn \\ &\cong 1 + \alpha^2 \frac{\underline{\delta}_r^2 + \bar{\tau}_p^2}{\underline{\delta}_r^2 + \alpha^2 \bar{\tau}_p^2} \left[\left(\frac{4N}{\pi} \right)^{1/2} - \left(\frac{8}{\pi} \right)^{1/2} \right]. \end{aligned} \quad (6.22)$$

In the particular case of perfect tuning of all resonant circuits $\underline{\delta}_r = 0$ and

$$(\underline{\Delta}_n/\underline{\Delta}_1)^2 \cong \left(\frac{1}{\pi n}\right)^{1/2}, \quad (6.23)$$

$$C^2 \cong 1 + \left(\frac{4N}{\pi}\right)^{1/2} - \left(\frac{8}{\pi}\right)^{1/2},$$

$$C \cong \left(\frac{4N}{\pi}\right)^{1/4}. \quad (6.24)$$

The last expression gives the factor by which the rms timing deviation at the output of repeater N is greater than at the output of the first repeater. The rms deviation at the output of the first repeater is greater than at the input by the factor α . The rms deviation at the output of repeater N is thus greater than at the input of the first repeater by the factor,

$$C_1 = \alpha \left(\frac{4N}{\pi}\right)^{1/4}. \quad (6.25)$$

For this particular case ($\delta_r = 0$) expressions equivalent to those above have been derived in unpublished work by H. E. Rowe of Bell Telephone Laboratories.

In accordance with (6.22) and (6.24) the cumulation of random timing deviations increases indefinitely with N when retiming is complete. The cumulation factor as given by (6.24) is in fact the same as would be obtained if a timing wave were transmitted on a separate pair, with a resonant circuit at each repeater to limit noise and with amplification of the timing wave at each repeater to obtain the same amplitude of the timing wave as when it is derived from the pulse train. With partial retiming cumulation is limited, for the reason that there is partial regeneration of both the pulse train and the timing wave.

Although with complete retiming the cumulation factor increases indefinitely with N , this is of but little practical significance, because of the slow rate of cumulation. At the output of a chain of N repeaters an rms deviation approximately equal to that at the input of the first repeater could be tolerated, in which case $C_1 \cong 1$. On this basis the permissible number of repeaters would be

$$N \cong \frac{\pi}{4} \frac{1}{\alpha^4} = \frac{\pi}{4} \left(\frac{Q}{\pi}\right)^2, \quad (6.26)$$

$$\cong 800 \quad \text{when} \quad Q = 100.$$

This assumes exact tuning of all resonant circuits. With mistuning of the resonant circuits, the permissible number of repeaters in tandem for a specified rms deviation at the output of the final repeater can be

determined with the aid of the cumulation factor given by (6.22). For example, if the rms deviation at the output of repeater N is assumed the same as at the input of the first repeater, the permissible number of repeaters in tandem is less than given by (6.26) by the factor $[(1 - m^2)/(1 + m^2)]^2$, $m = \delta_r/\bar{\tau}_p$. When the fixed phase shift is 30° , $m \cong 0.5$ and $N \cong 300$.

6.5 Self-Starting of Self-Timed Repeaters

With self-timing it is necessary that repeaters be self-starting if the timing wave should be absent for any reason. If each repeater is self-starting, this will also be the case for a repeater chain, since starting will be progressive along the chain. Initially, before the timing wave has reached the appropriate amplitude at all repeaters, there will be a high rate of digital errors.

With timing from the received pulse train, the resonant circuit will be excited by every pulse and the timing wave will reach its normal amplitude in about $n \cong Q$ pulses. With timing from the regenerated pulse

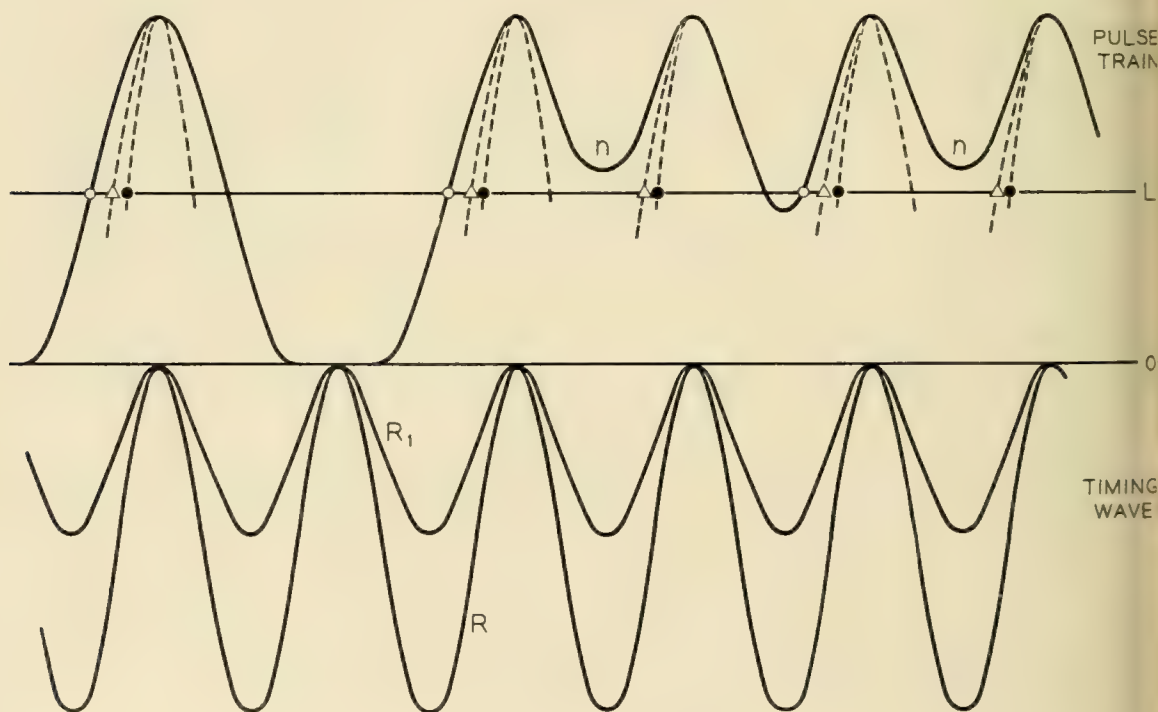


Fig. 9 — Progression of repeater starting in absence of timing wave when timing is derived from regenerated pulse train.

- Triggering points with timing wave absent. Noise prevents triggering at certain points, n . Timing wave reaches fraction of normal value, R_1 .
- △ Triggering points with timing wave R_1 . Timing wave increases to normal amplitude R .
- Triggering points with normal timing wave.

train the resonant circuit will not be excited by every pulse, unless the shape of the received pulses is such that there are virtually no overlaps between pulses so that the triggering level will be penetrated by each pulse.

With a pulse shape as assumed in the previous analysis, the amplitude of a pulse train midway between pulses is half the peak amplitude of the pulses, as indicated in Fig. 9. In the presence of noise, triggering will in this case occur on the average for every second pulse, as indicated in the above figure. If it is assumed that the resonant circuit has the maximum permissible phase shift of about 20° allowed with timing from the output, the amplitude of the timing wave with excitation from every pulse will be virtually equal to the peak pulse amplitude. With excitation from half the pulses, the amplitude of the timing wave will rapidly reach half the peak amplitude of the pulses. When this initial timing wave is combined with the pulse train, triggering will occur for virtually all pulses, as indicated in Fig. 9. It will thus reach its normal value. If the phase shift is greater than 20° as assumed above, say 60° , the initial amplitude of the timing wave will be $\frac{1}{4}$ the peak pulse amplitude. Combination of this initial timing wave with the pulse train will increase the number of pulses exciting the resonant circuit, which in turn increases the amplitude of the timing waves, etc.

Self-starting with a pulse shape as assumed in this analysis is thus insured.

VII SUMMARY

In self-timing regenerative repeaters as considered here, a timing wave is derived from either the received or regenerated pulse train with the aid of a simple resonant circuit tuned to the pulse repetition frequency. This timing wave is combined linearly with received pulse trains as indicated in Fig. 1, and pulses are regenerated when the combined wave penetrates a certain triggering level.

It is concluded that if these timing principles are implemented by appropriate repeater instrumentation, a performance can be realized that approaches that of ideal regenerative repeaters. To this end it is necessary to meet certain requirements with regard to the loss constant Q of the resonant circuit, its frequency precision, the shape of received pulses and the amplitude of the timing wave in relation to that of received pulses.

Equalization of each repeater section should preferably be such that the received pulses have a shape and duration in relation to the pulse

interval as indicated in Fig. 3, and the peak amplitude of the timing wave should be about equal to that of the received pulses. Under these conditions the pulse repetition frequency will be present in the received pulse train in sufficient amplitude to permit derivation of the timing wave from the received pulse train, and to permit rapid self-starting in the absence of a timing wave if it is derived from the regenerated pulse train.

A loss constant of the resonant circuit $Q \cong 100$ appears desirable. This value is sufficiently low to be readily realized with simple resonant circuits consisting of a coil and capacitor in series or parallel, without unduly severe requirements on its frequency precision. It is also adequately high from the standpoint of avoiding excessive random timing deviations in regenerated pulses from amplitude and phase deviations in the timing wave.

The tolerable deviation in the resonant frequency from the pulse repetition frequency with $Q = 100$ is about 0.2 per cent when the timing wave is derived from the received pulse train, and about 0.06 per cent when it is derived from the regenerated pulse train. These frequency precisions correspond to a maximum fixed phase shift of 15° in the timing wave, and allow for the possibility that random timing deviations resulting from amplitude variations in the timing wave may cumulate directly along a repeater chain, rather than on a root-sum-square basis. With root-sum-square cumulation of timing deviations from all sources, the frequency deviations could be about twice as great.

When the above requirements are met the reduction in the tolerance to noise owing to timing deviations in a repeater chain is limited to about 2 db. If the requirements on frequency precision of the resonant circuit are met, substantial degradation or improvement in performance would not be expected as a result of moderate changes in the other design parameters.

VIII ACKNOWLEDGMENTS

In this presentation the writer had the benefit of unpublished work, referred to previously, by W. R. Bennett and J. R. Pierce on the derivation of a timing wave from a pulse train with the aid of a resonant circuit, and by H. E. Rowe on the cumulation of timing deviations in a chain of repeaters with complete retiming. Bennett, Pierce and Rowe are at Bell Telephone Laboratories. He is also indebted to H. E. Rowe for a critical review that resulted in several improvements in the analysis.

APPENDIX

IX RESONANT CIRCUIT RESPONSE TO RANDOM BINARY PULSE TRAINS

1 General

In the following analysis of the response of a resonant circuit to a binary "on-off" pulse train, the pulses are assumed of sufficiently short duration to be regarded as impulses. This is a legitimate approximation when the duration does not exceed about half the interval between pulses.

The pulse train is regarded as made up of three components, as indicated in Fig. 10. The first is a systematic component consisting of pulses of amplitude $\frac{1}{2}$. This component gives rise to a steady state response at the fundamental frequency of the pulse sequence. The second component consists of pulses of amplitude $\pm\frac{1}{2}$, with random $\pm$ polarity. This component gives rise to a random component in the resonant cir-

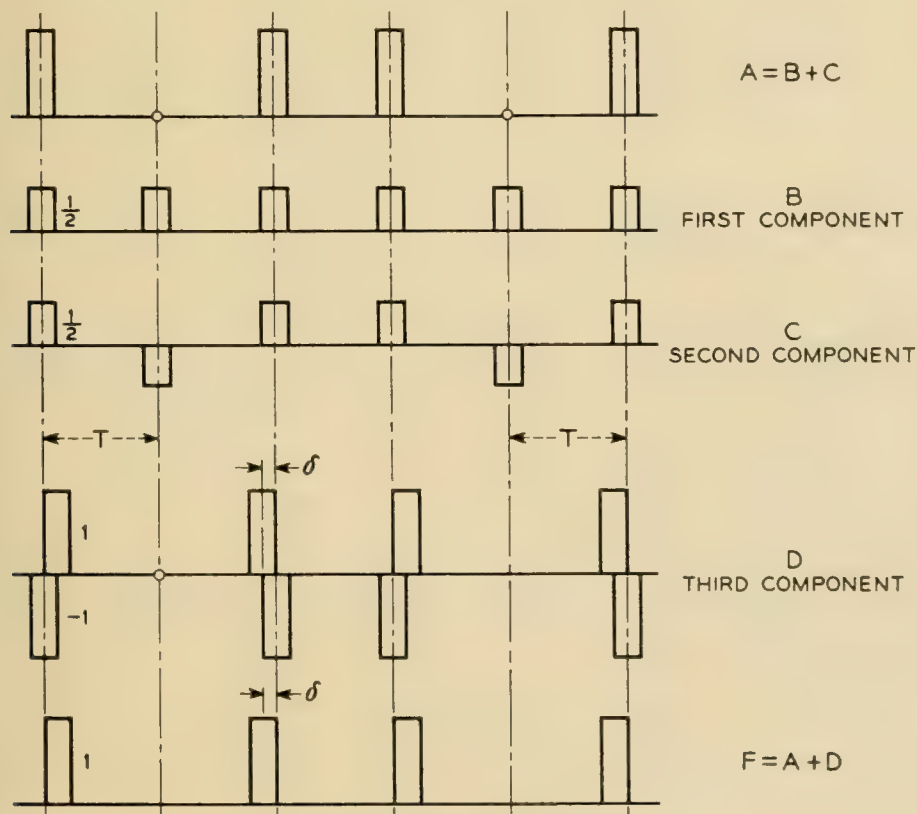


Fig. 10 — Components of random binary on-off pulse train. A. — Transmitted "on-off" pulses. B. — Steady state pulse train of fundamental frequency $f = 1/T$. C. — Random pulse train with zero mean value. D. — Random pulse train with displacements $\pm\delta$. F. — "On-off" pulse train with displacements $\pm\delta$ from average pulse interval T .

cuit response; i.e., a fluctuation about the steady state value derived from the first component.

The third component consists of a train of dipulses. Each dipulse consists of a pair of pulses of amplitude 1 and -1 , displaced by an interval $\pm\delta$. The response of the resonant circuit to this component gives the effect of random displacements $\pm\delta$ in the original "on-off" pulse train.

2 Impedance of Resonant Circuit

The impedance of a resonant circuit consisting of R , L and C in parallel is

$$Z(i\omega) = \bar{Z}(i\omega)e^{-i\psi}, \quad (1)$$

$$\bar{Z}(i\omega) = \frac{R}{[1 + Q^2(\omega/\omega_0 - \omega_0/\omega)^2]^{1/2}} = R \cos \psi, \quad (2)$$

$$\tan \psi = Q(\omega/\omega_0 - \omega_0/\omega), \quad (3)$$

where

$$Q = \omega_0 RC = \text{Loss constant}, \quad (4)$$

and

$$\omega_0 = (1/LC)^{1/2} = \text{Resonant frequency}. \quad (5)$$

The above expressions also apply for the admittance of a resonant circuit consisting of R , L and C in series, except that in this case $Q = \omega_0 L/R$.

3 Impulse Response of Resonant Circuit

When a rectangular pulse of unit amplitude and sufficiently short duration δ_0 is applied to a resonant circuit, the impulse response for $Q \gg 1$ is of the form

$$P(t) = P(0) \cos \omega_0 t e^{-\omega_0 t/2Q}, \quad (6)$$

where

$$P(0) = \omega_0 \delta_0 R/Q. \quad (7)$$

$P(t)$ designates voltage in response to an impulse current in the case of a parallel resonant circuit, or the current in response to an impulse voltage in the case of a series resonant circuit.

4 Response to Steady State Impulse Train

Let a long sequence of impulses of amplitude $\frac{1}{2}$ and the same polarity impinge on a resonant circuit at uniform intervals T . The response after N impulses is then

$$A_s(t) = \frac{1}{2} \sum_{n=0}^N P(t - nT) \quad (8)$$

$$= \frac{P(0)}{2} \sum_{n=0}^N \cos \omega_0(t - nT) e^{-\omega_0(t-nT)/2Q}. \quad (9)$$

The subscript s indicates a systematic component.

The above series is conveniently summed by taking the real part of the series

$$A_s'(t) = \frac{P(0)}{2} \sum_{n=0}^N e^{i\omega_0(t-nT)} e^{-\omega_0(t-nT)/2Q}. \quad (10)$$

With $t = NT + t_0$, $0 < t_0 < T$:

$$A_s'(t) = \frac{P(0)}{2} e^{i\omega_0 t_0} e^{-\omega_0 t_0/2Q} \sum_{n=0}^N e^{i\omega_0 t(N-n)} e^{-\omega_0 t(N-n)/2Q}. \quad (11)$$

When $N \rightarrow \infty$, the steady state responses becomes

$$A_s'(t) = \frac{P(0)}{2} e^{i\omega_0 t_0} e^{-\omega_0 t_0/2Q} \frac{1}{1 - e^{i\omega_0 T} e^{-\omega_0 T/2Q}}. \quad (12)$$

The interval between pulses can be written

$$T = 2\pi/\omega, \quad (13)$$

where ω is the fundamental frequency of the impulse train, or the pulse repetition frequency.

Let

$$\omega_0 = \omega - \Delta\omega,$$

so that

$$\omega_0 = \frac{2\pi}{T} (1 - \Delta\omega/\omega). \quad (14)$$

The following approximations then apply:

$$e^{+i\omega_0 T} = e^{2\pi i} e^{-2\pi i \Delta\omega/\omega} = e^{-2\pi i \Delta\omega/\omega}, \quad (15)$$

$$\cong 1 - 2\pi i \Delta\omega/\omega;$$

$$e^{-\omega_0 T/2Q} = e^{-\pi/Q} e^{+(\pi/Q) \Delta\omega/\omega}, \quad (16)$$

$$\cong 1 - \pi/Q \text{ when } \pi/Q \ll 1.$$

With these approximations

$$1 - e^{i\omega_0 T} e^{-\omega_0 T/2Q} \cong \frac{\pi}{Q} [1 + i\psi], \quad (17)$$

where

$$\psi = \frac{2\Delta\omega}{\omega} Q, \quad (18)$$

will be recognized as the phase shift of the resonant current at the frequency ω , as obtained from (3) with $\omega = \omega_0 + \Delta\omega$.

A further approximation that can be introduced in (12) is

$$\begin{aligned} e^{i\omega_0 t_0} e^{-\omega_0 t_0/2Q} &= e^{i\omega t_0} e^{-i\Delta\omega t_0} e^{-\omega_0 t_0/2Q}, \\ &\cong e^{i\omega t_0}, \end{aligned} \quad (19)$$

since $t_0 < T$ and $\Delta\omega t_0$ and $\omega_0 t_0/2Q \ll 1$.

With the above approximations (12) becomes

$$A_s' = \frac{P(0)}{2} \frac{Q}{\pi} e^{i[\omega t_0 - \psi]} \cos \psi. \quad (20)$$

The real part of this expression is

$$A_s = \frac{P(0)}{2} \frac{Q}{\pi} \cos (\omega t_0 - \psi) \cos \psi, \quad (21)$$

which is the response to the steady state component of the pulse train.

5 Response to Random Component of Impulse Train

Let a sequence of impulses of amplitude $\frac{1}{2}$ and randomly positive and negative polarity impinge on the resonant circuit at intervals T . The response is then,

$$A_r(t) = \frac{P(0)}{2} \sum_{n=0}^N \pm \cos \omega_0(t - nT) e^{-\omega_0(t-nT)/2Q}. \quad (22)$$

This expression for the random component differs from (9) for the systematic component in that the impulses have random $\pm$ polarity. If all signs are chosen the same, the values of $A_r(t)$ will be either $-A_s(t)$ or $+A_s(t)$. The resultant response of the resonant circuit, i.e. $A_s(t) + A_r(t)$, can thus vary between the limit 0 and $2A_s(t)$. $A_r(t)$ represents a random fluctuation about $A_s(t)$ as a mean value. In the following the rms value of this fluctuation is evaluated.

In order to determine the components of $A_r(t)$ in phase and at quadra-

ture with the steady state response as given by (21), it is convenient to write

$$\omega_0 = \omega - \Delta\omega,$$

$$\begin{aligned}\cos \omega_0(t - nT) &= \cos [\omega(t - nT) - \psi + \psi - \Delta\omega(t - nT)] \\ &= \cos [\omega(t - nT) - \psi] \cos [\psi - \Delta\omega(t - nT)] \\ &\quad - \sin [\omega(t - nT) - \psi] \sin [\psi - \Delta\omega(t - nT)].\end{aligned}\quad (23)$$

With $t = NT + t_0$, and $\omega T = \pi$, (22) can be written:

$$\begin{aligned}A_r(t) &= \cos(\omega t_0 - \psi) \sum_{n=0}^N \pm \cos[\psi_1 - \Delta\omega T(N - n)] e^{-\omega_0(T)(N-n)/2Q} \\ &\quad - \sin(\omega t_0 - \psi) \sum_{n=0}^N \pm \sin[\psi_1 - \Delta\omega T(N - n)] e^{-\omega_0(T)(N-n)/2Q}\end{aligned}\quad (24)$$

where $\psi_1 = \psi - \Delta\omega t_0 = \psi \left(1 - \frac{\omega t_0}{2Q}\right) \cong \psi$, since $\omega t_0/2Q \leq \pi/2Q \ll 1$.

With equal probabilities of a plus or a minus sign in the summations, the rms value of the in-phase component becomes

$$\begin{aligned}\underline{A}_r' &= \left[\sum_{n=0}^N \cos^2[\psi - \Delta\omega T(N - n)] e^{-\omega_0 T(N-n)/Q} \right]^{1/2} \\ &= \left[\sum_{n=0}^N \frac{1}{2} (1 + \cos 2[\psi - \Delta\omega T(N - n)] e^{-\omega_0 T(N-n)/Q}) \right]^{1/2}.\end{aligned}\quad (25)$$

The rms value of the quadrature component becomes

$$\begin{aligned}\underline{A}_r'' &= \left[\sum_{n=0}^N \sin^2[\psi - \Delta\omega T(N - n)] e^{-\omega_0 T(N-n)/Q} \right]^{1/2} \\ &= \left[\sum_{n=0}^N \frac{1}{2} (1 - \cos 2[\psi - \Delta\omega T(N - n)] e^{-\omega_0 T(N-n)/Q}) \right]^{1/2}.\end{aligned}\quad (26)$$

These expressions can be transformed into sums of geometric series by writing

$$\cos x = \frac{1}{2}(e^{ix} + e^{-ix}), \quad x = 2[\psi - \Delta\omega T(N - n)].$$

Evaluation of (25) and (26) by this method gives

$$\underline{A}_r' = \frac{P(0)}{2} \frac{1}{2^{1/2}} \left[\frac{1}{1 - e^{-\omega_0 T/Q}} + \frac{N}{D} \right]^{1/2}, \quad (27)$$

$$\underline{A}_r'' = \frac{P(0)}{(2)} \frac{1}{2^{1/2}} \left[\frac{1}{1 - e^{-\omega_0 T/Q}} - \frac{N}{D} \right]^{1/2}, \quad (28)$$

where

$$N = \cos 2\psi(1 - \cos 2\Delta\omega T e^{-\omega_0 T/2Q}) + \sin 2\psi \sin 2\Delta\omega T e^{-\omega_0 T/2Q}, \quad (29)$$

$$D = 1 + e^{-2\omega_0 T/Q} - 2e^{-\omega_0 T/Q} \cos 2\Delta\omega T. \quad (30)$$

With the same approximations as used previously in connection with (12) and with

$$\cos 2\Delta\omega T \cong 1 - 2(\Delta\omega T)^2, \quad \sin 2\Delta\omega T \cong 2\Delta\omega T,$$

$$N \cong \frac{2\pi}{Q}, \quad (31)$$

$$D \cong \left(\frac{2\pi}{Q}\right)^2 [1 + \psi^2], \quad (32)$$

$$1 - e^{-\omega_0 T/Q} = 2\pi/Q. \quad (33)$$

With these approximations in (27) and (28),

$$\underline{A}_r' \cong \frac{P(0)}{2} \left(\frac{Q}{2\pi}\right)^{1/2} [1 - \psi^2/2]^{1/2}, \quad (34)$$

$$\underline{A}_r'' \cong \frac{P(0)}{2} \left(\frac{Q}{2\pi}\right)^{1/2} \frac{|\psi|}{2^{1/2}}, \quad (35)$$

which apply when ψ is small and $(2\pi/Q) \ll 1$.

6. Response to Random Dipulse Train

Each dipulse is assumed to consist of two impulses of unit amplitude and opposite polarity, displaced by an interval δ , which in general will be a function of the pulse position; i.e., $\delta = \delta(n)$. The response of the resonant circuit to a train of such dipulses, obtained by taking the difference in response to the two impulses, is given by

$$A_\delta(t) = P(0) \left[\sum_{n=0}^N \cos \omega_0(t - nT) e^{-\omega_0(t-nT)/2Q} - \cos \omega_0[t - nT + \delta(n)] e^{-\omega_0[t-nT+\delta(n)]/2Q} \right]. \quad (36)$$

In determining the response, mistuning of the resonant current can be disregarded; i.e., $\omega_0 = \omega$. Furthermore, in the second term of (36) it is permissible to take

$$\exp [-\omega_0 \delta(n)/2Q] \cong 1.$$

With the following further approximations

$$\begin{aligned}\cos \omega_0(t - nT) - \cos \omega_0[t - nT + \delta(n)] \\ &= \sin \omega_0[t + \delta(n)/2] 2 \sin [\omega_0\delta(n)/2], \\ &\cong \omega_0\delta(n) \sin \omega_0 t,\end{aligned}\quad (37)$$

expression (36) becomes:

$$\begin{aligned}A_\delta(t) &= P(0)\omega_0 \sin \omega_0 t \sum_{n=0}^N \delta(n) e^{-\omega_0(t-nT)/2Q} \\ &= P(0)\omega_0 \sin \omega_0 t_0 \sum_{n=0}^N \delta(n) e^{-\omega_0 T(N-n)/2Q},\end{aligned}\quad (38)$$

where the substitution $t = NT + t_0$ has been made as in previous expressions.

The above expression shows that the resonant circuit response will be at quadrature with the steady state timing wave $\cos \omega t_0$.

In the above expressions, the dipulses are assumed to be present at intervals T , whereas in a random pulse train they will be present at average intervals $2T$. The rms value of the quadrature component with randomly positive and negative dipulses at intervals $2T$, with an rms displacement $\hat{\delta}$, is

$$\begin{aligned}\underline{A''}_\delta &= P(0)\omega_0\hat{\delta} \left[\sum_{n=0}^N e^{-2\omega_0 T(N-n)/Q} \right]^{1/2} \\ &= \frac{P(0)}{2} \omega_0\hat{\delta} \left(\frac{Q}{\pi} \right)^{1/2}.\end{aligned}\quad (39)$$

In (38) the function $e^{-\omega_0 t/2Q}$ will be recognized as the impulse response function of a circuit with impedance

$$Z(i\omega) = \frac{1}{\beta + i\omega}, \quad \beta = \omega_0/2Q, \quad (40)$$

$$= \bar{Z}(i\omega) e^{-i\psi}, \quad (41)$$

$$\bar{Z}(i\omega) = \frac{1}{\beta} \left[\frac{1}{1 + \omega^2/\beta^2} \right]^{1/2}, \quad (42)$$

$$\tan \psi = \omega/\beta. \quad (43)$$

It will also be recognized that (39) corresponds to the rms response of such a circuit, when impulses $\delta(n)$ of random amplitude with an rms value $\hat{\delta}$ are applied to average intervals $2T$. Thus (39) can alternately be obtained from

$$\begin{aligned}\underline{A}_{\delta}'' &= P(0)\omega_0\hat{\omega} \left[\frac{1}{2T} \frac{1}{2\pi} \int_{-\infty}^{\infty} [\bar{Z}(i\omega)]^2 d\omega \right]^{1/2} \\ &= P(0)\omega_0\hat{\omega} \left[\frac{1}{4\pi T\beta^2} \beta(\tan^{-1}\omega/\beta)_{-\infty}^{\infty} \right]^{1/2}\end{aligned}\quad (44)$$

$$\begin{aligned}&= P(0)\omega_0\hat{\omega} \left(\frac{1}{4T\beta} \right)^{1/2} \\ &= \frac{P(0)}{2} \omega_0\hat{\omega} \left(\frac{Q}{\pi} \right)^{1/2}.\end{aligned}\quad (45)$$

Let the output of the first resonant circuit be applied to a second resonant circuit, and in turn to n successive resonant circuits, with an amplitude amplification β between successive resonant circuits. At the output of the n^{th} resonant circuit, the rms amplitude of the response is then obtained from

$$\begin{aligned}\underline{A}_{\delta,n}'' &= P(0)\omega_0\hat{\omega} \left[\frac{\beta^{2(n-1)}}{4\pi T} \int_{-\infty}^{\infty} [\bar{Z}^2(i\omega)]^n d\omega \right]^{1/2} \\ &= P(0)\omega_0\hat{\omega} \left[\frac{1}{4\pi T\beta^2} \int_{-\infty}^{\infty} \frac{d\omega}{(1 + \omega^2/\beta^2)^n} \right]^{1/2} \\ &= \frac{P(0)}{2} \omega_0\hat{\omega} \left(\frac{Q}{\pi} \right)^{1/2}, \quad I_n = \underline{A}_{\delta}'' I_n,\end{aligned}\quad (46)$$

where

$$I_n^2 = \frac{1}{\pi\beta} \int_{-\infty}^{\infty} \frac{d\omega}{(1 + \omega^2/\beta^2)^n}, \quad (47)$$

$$= 1, \quad n = 1,$$

$$= \frac{2n-3}{2(n-1)} I_{n-1}^2, \quad n \geq 2, \quad (48)$$

$$= I_{n-1}^2 \left(1 - \frac{1}{2(n-1)} \right),$$

$$I_2^2 = (1 - \frac{1}{2}), \quad I_3^2 = (1 - \frac{1}{4})I_2^2, \quad I_4^2 = (1 - \frac{1}{6})I_3^2.$$

Thus (46) can be written:

$$\underline{A}_{\delta,n}'' = \underline{A}_{\delta}'' \alpha_2 \alpha_3 \cdots \alpha_n, \quad (49)$$

where

$$\alpha_j^2 = 1 - \frac{1}{2(j-1)}, \quad (50)$$

$$\alpha_2^2 \alpha_3^2 \cdots \alpha_n^2 = \left(1 - \frac{1}{2}\right) \left(1 - \frac{1}{4}\right) \left(1 - \frac{1}{6}\right) \cdots \left(1 - \frac{1}{2(n-1)}\right) \quad (51)$$

$$= \frac{1 \cdot 3 \cdot 5 \cdot 7 \cdots [2(n-1) - 1]}{2 \cdot 4 \cdot 6 \cdot 8 \cdots 2(n-1)} \quad (52)$$

$$= \frac{(2n)!}{2^{2n}(n!)^2}. \quad (53)$$

When $n \gg 1$, (51) approaches the value

$$\alpha_2^2 \alpha_3^2 \cdots \alpha_n^2 \cong \left(\frac{1}{\pi n}\right)^{1/2}. \quad (54)$$

The latter approximation is based on the following expression, for $x = -\frac{1}{2}$, given in Whittaker and Watson's: "Modern Analysis" page 259:

$$\lim_{n \rightarrow \infty} (1+x)(1+x/2)(1+x/3) \cdots (1+x/n) = \frac{n^x}{\Gamma(1+x)}, \quad (55)$$

where Γ is the gamma function, $\Gamma(-\frac{1}{2} + 1) = \pi^{1/2}$.

The above analysis assumes that the timing wave at each resonant circuit is applied directly to the next resonant circuit, except for the amplification between resonant circuits. This would be the case if the timing wave were transmitted on a separate pair, in which case $A''_{\delta,n}$ would be the rms quadrature component owing to noise in the timing circuit.

In regenerative repeaters, deviations in the timing wave resulting from the quadrature component are imparted at intervals T into the next repeater section as deviations in the spacing of pulses. These timing deviations occurring at intervals T will have a certain random amplitude distribution, which can be regarded as having a certain frequency spectrum. When the deviations are discrete and occur at intervals T , the spectrum will extend to a maximum frequency $f_{\max} = 1/2T$, or $\omega_{\max} = \pi/T = \omega_0/2$. In this case the upper and lower limits of the integrals above would be replaced by $\pm \omega_0/2$, except for the first repeater section. The recurrence relation (48) is then no longer exact, but the resultant modification is insignificant and can be disregarded. This will be seen when the value $\omega_0/2$ is inserted for ω in the integrand of (47), which then becomes $1/(1 + Q^2)$, as compared with 1 for $\omega = 0$. Thus the contribution to the integrals for $\omega > \omega_0/2$ can for practical purposes be disregarded.

A Sufficient Set of Statistics for a Simple Telephone Exchange Model

By V. E. BENEŠ

(Manuscript received October 17, 1956)

This paper considers a simple telephone exchange model which has an infinite number of trunks and in which the traffic depends on two parameters, the calling-rate and the mean holding-time. It is desired to estimate these parameters by observing the model continuously during a finite interval, and noting the calling-time and hang-up time of each call, insofar as these times fall within the interval. It is shown that the resulting information may, for the purpose of this estimate, be reduced without loss to four statistics. These statistics are the number of calls found at the start of observation, the number of calls arriving during observation, the number of calls terminated during observation, and the average number of calls existing during the interval of observation. The joint distribution of these sufficient statistics is determined, in principle, by deriving a generating function for it. From this generating function the means, variances, covariances, and correlation coefficients are obtained. Various estimators for the parameters of the model are compared, and some of their distributions, means, and variances presented.

I THEORETICAL PROBLEMS AND METHODS OF TRAFFIC MEASUREMENT

Four important kinds of theoretical problems arise in the measurement of telephone traffic. These are: (1) the choice of a mathematical model, containing parameters characteristic of the traffic, to serve as a description; (2) the devising of efficient methods of estimating the parameters; (3) the determination of the anticipated accuracy of measurements; and (4) the assessment of actual accuracy, after measurements have been made.

The present paper deals with aspects of the second and third kinds of problem, for the simplest and least realistic mathematical model of telephone traffic. Specifically, for this model, we treat the problems of (i) complete extraction of the information from a given observation period,

without regard to costs of observation, and (ii) determination of the anticipated accuracy of certain methods of estimation which arise naturally from the discussion of complete extraction.

The method by which we attack problems (i) and (ii) in this paper has three stages. First we choose a small number of significant properties of, or factors in, the physical system we are studying. Then we abstract these properties into a mathematical model of the physical system. Finally, from the properties of the model, we derive results which may be interpreted as answers to the two problems treated. The advantage of this method is that we can use the precise, powerful apparatus of mathematics in studying the model; its limitation is that it yields results which are only as accurate as the model in describing reality.

A method similar to the above forms the theoretical underpinning of telephone traffic engineering itself. To design equipment effectively, the traffic engineer needs a description of the traffic that is handled by central offices. He decides what properties of the entire system of telephone equipment and customers will be most useful to him in describing the traffic. He then designates certain parameters to serve as mathematically precise idealizations of these properties, and in terms of these parameters constructs a model of the traffic, upon which he bases much of his engineering.

In choosing a mathematical model for a physical system, one is confronted with two generally opposed desiderata: fidelity to the system described, and mathematical simplicity. The model may involve important departures from physical reality; a model that is sufficiently amenable to mathematical analysis often results only after one has introduced admittedly false assumptions, ignored certain effects and correlations, and generally oversimplified the system to be studied. However, the abstract model will be an exact and simple tool for analysis.

We can construct a simple mathematical model for the operation of a telephone central office by leaving out of consideration many important facts about such systems, and by concentrating on factors most relevant to operation. Since we are interested in telephone traffic and in the availability of plant, it seems natural to require that a realistic model take account of at least the following five significant factors: (1) the demand for telephone service; (2) the rate at which requests for service can be processed and connections established; (3) the lengths of conversations; (4) the supply of central office equipment; and (5) the manner in which the first four factors are interrelated. Unfortunately, the mathematical complexity of such a realistic model precludes easy investigation. Therefore, the model used in this paper is based only on factors (1) and (3).

The demand for telephone traffic is usually made precise by describing a stochastic process which represents the way in which requests for telephone service occur in time. A realistic description will take account of the facts that, the demand is not constant, but has daily extremes, and that in small systems, the demand may be materially lessened when many conversations are in progress. Since taking account of the first fact leads to a more complicated model in which our investigations are more difficult, we ignore it, with the proviso that the results we derive are only applicable to systems and observations for which the demand is nearly constant. The second kind of variation in demand becomes insignificant as the number of subscribers increases and the traffic remains constant. Hence, we further confine the applicability of our results to systems with large numbers of subscribers, and we assume that the demand does not depend on the number of conversations in existence.

With these assumptions, a mathematically convenient description of the demand is specified by the condition that the time-intervals between requests for service have lengths which are mutually independent positive random variables, with a negative exponential distribution.

A telephone central office contains two kinds of equipment: control circuits which establish a desired connection, and talking paths over which a conversation takes place. The time that a request for service occupies a unit of equipment, be the unit a control circuit or a talking path, is called the holding-time of the unit. A request for service affects the availability of both kinds of equipment but, except for special cases, the holding-times of talking paths are usually much longer than the holding-times of control units such as markers, connectors, or registers. In view of this disparity, we assume that the only holding-times of consequence are the lengths of conversations; i.e., the holding-times of talking paths. We assume also that these lengths are mutually independent positive random variables, with a negative exponential distribution.

For the simplest mathematical model of telephone traffic, we may consider the arrangement of switches and transmission lines which constitutes a talking path in the physical office to be replaced by an abstract unit called a "trunk". A trunk is then an abstraction of the equipment made unavailable by one conversation, and we may measure the supply of talking paths in the office by the number of trunks in a model. The word "trunk" is also used to mean a transmission line linking two central offices, but as long as we have explained our use of the word there need be no confusion. Often the number of transmission lines leading out of an office is a major limitation on its capacity to carry conversations, and in this case the two uses of the word "trunk" are very similar. Un-

fortunately, we do not take advantage of this similarity, since we make the mathematically convenient but wholly unrealistic assumption that the number of trunks in the model is infinite.

The model we investigate thus depends on only two of the factors previously listed as essential to a realistic model: namely, (1) the demand for service, and (3) the lengths of conversations. In view of the simplicity and inaccuracy of this model, the question arises whether much is gained from a detailed analysis. Such scrutiny may indeed reveal little that is of great practical value to traffic engineers. It is important methodologically, however, to have a detailed treatment of at least one approximate case. We undertake this detailed treatment largely for the insight that it may give into methods which could be useful in dealing with more complex and more accurate models.

Once a designer has chosen a model and has specified the parameters he would like to have measured, it is up to the statistician to invent efficient means of measurement, by choosing, for each parameter, some function of possible observations to serve as an estimate of that parameter. One measure of efficiency that is of mostly theoretical interest is the observation time required to achieve a given degree of anticipated accuracy; the most realistic measure of efficiency is in terms of dollars and man-hours. It may often be more efficient, in the sense of the latter measure, to spread observation over enough more time to compensate for the inability of an intrinsically cheaper method of measurement to extract all of the information present in a fixed time of observation. For example, periodic scanning of switches in a telephone exchange is usually less costly than continuous observation. As a result, telephone traffic measurement is usually carried out by averaging sequences of instantaneous periodic observations of the number of calls present, rather than by continuous time averaging, although it can be shown that continuous observation is more efficient at extracting information. Thus statistical efficiency, which may be expensive in terms of measuring equipment, can be exchanged for observation time, which may be less costly. This exchange brings about a reduction in cost without impairing accuracy.

Our concern in this paper is with the less practical problems of complete extraction, and of the anticipated accuracy of estimation methods based on complete extraction. Let us consider how our mathematical model can shed light on these problems. A mathematical model may or may not be a faithful description of the behavior of real telephone systems. Nevertheless random numbers, with or without modern computing machines, enable one to make experiments and observations on physical situations which approximate, arbitrarily closely, any mathematical model. Thus we can speak meaningfully of events in the model, and of

making measurements and observations on the model. The mathematical model elucidates our problems in the following ways: (1) it enables us to state precisely what information is provided by observation; (2) it enables us to explain what we mean by complete extraction of information; and (3) it enables us to derive results about the anticipated accuracy of measurements in the model. These results will have approximately true analogues in physical situations to which the model is applicable.

The calls existing during the observation interval (O, T) fall into four categories: (i) those which exist at O , and terminate before T ; (ii) those which fall entirely within (O, T) ; (iii) those which exist at O and last beyond T ; and (iv) those which begin within (O, T) and last beyond T . For calls of category (i), we assume that we observe the hang-up time of each call; for category (ii), we observe the matching calling-time and hang-up time of each conversation; for category (iii), we observe simply the number of such calls; and for category (iv), we observe the calling-times. Table I summarizes the kinds of calls and the information observed about each.

What we mean by the complete extraction of information is made precise by the statistical concept of *sufficiency*. By a statistic we mean any function of the observations, and by an estimator we mean a statistic which has been chosen to serve as an estimate of a particular parameter. Roughly and generally, a set S of statistics is sufficient for a set P of parameters when S contains all the information in the original data that was relevant to parameters in P . If S is sufficient for P , there is a set E of estimators for parameters in P , such that the estimators in E depend only on statistics from S , and such that an estimator from E does at least as well as any other estimator we might choose for the same parameter. Thus we incur no loss in reducing the original data (of specified form) to the set S of statistics. It remains to state what it means for S to contain all the relevant information. We do this in terms of our model.

The mathematical model we are adopting contains two distribution

TABLE I — INFORMATION OBSERVED

Types of Calls	Start in (O, T)	Start before O
End in (O, T)	(ii), matching calling-times and hang-up times known, number of calls known	(i), hang-up times known, number of calls known
End after T	(iv), calling-times known, number of calls known	(iii), number of calls known

functions, that of the intervals between demands for service, and that of the lengths of conversations. We have supposed that these distributions are both of negative exponential type, each depending on a single parameter. Thus we know the functional form of each distribution, and each such form has one unknown constant in it. Since the mathematical structure of the model is fully specified except for the values of the two unknown constants, we can assign a likelihood or a probability density to any sequence Σ of events in the model during the interval (O, T) . This likelihood will depend on the parameters, on Σ , and on the number of calls in existence at the start O of the interval. If the likelihood $L(\Sigma)$ can be factored into the form $L = F \cdot H$, where F depends on the parameters and on statistics from the set S only, and H is independent of the parameters, then the set S of statistics may be said to summarize all the information (in a sequence Σ) relevant to the parameters. If L can be so factored, then S is sufficient for the estimation of the parameters.

The mathematical model to be used in this paper is described and discussed in Sections II and III, respectively. Section IV contains a summary of notations and abbreviations which have been used to simplify formulas.

In Appendix A we show that the original data we have allowed ourselves can be replaced by four statistics, which are sufficient for estimation. In Appendix B and Sections V–VIII we discuss various estimators (for parameters of the model) based on these four statistics. To determine the anticipated accuracy of these methods of measurement, we consider the statistics themselves as random variables whose distributions are to be deduced from the structure of the model.

A primary task is the determination of the joint distribution of the sufficient statistics. In view of the sufficiency, this joint distribution tells us, in principle, just what it is possible to learn from a sample of length T in this simple model. By analyzing this distribution we can derive results about the anticipated accuracy of measurements in the model.

The joint distribution of the sufficient statistics is obtainable in principle from a generating function computed in Appendix C, using methods exemplified in Section X. This generating function is the basic result of this paper. The implications of this result are summarized in Section IX, which quotes the generating function itself, and presents some statistical properties of the sufficient statistics in the form of four tables: (i) a table of generating functions obtainable from the basic one; (ii) a table of mean values; (iii) a table of variances and covariances; and (iv), a table of squared correlation coefficients. (The coefficients are all non-negative.)

II DESCRIPTION OF THE MATHEMATICAL MODEL

Throughout the rest of the paper we follow a simplified form of the notational conventions of J. Riordan's paper¹¹ wherever possible. A summary of notations is given in Section IV. The model we study has the following properties:

(i) Demands for service arise individually and collectively at random at the rate of a calls per second. Thus the chance of one or more demands in a small time-interval Δt is

$$a\Delta t + o(\Delta t),$$

where $o(\Delta t)$ denotes a quantity of order smaller than Δt . The chance of more than one demand in Δt is of order smaller than Δt . It can be shown (Feller,² p. 364 et seq.) that this description of the demand is equivalent to saying that the intervals between successive demands for service are all independent, with the negative exponential distribution

$$1 - e^{-at}.$$

This again is equivalent to saying that the call arrivals form a Poisson process;² i.e., that for any time interval, t , the probability that exactly n demands are registered in t is

$$\frac{e^{-at}(at)^n}{n!}.$$

Thus the number of demands in t has a Poisson distribution with mean at .

(ii) The holding-times of distinct conversations are independent variates having the negative exponential distribution

$$1 - e^{-\gamma t},$$

where γ is the reciprocal of the mean holding-time h . This description of the holding-time distribution is the same as saying that the probability that a conversation, which is in progress, ends during a small time-interval Δt is

$$\gamma\Delta t + o(\Delta t),$$

without regard to the length of time that the conversation has lasted (Feller, p. 375).

(iii) The model contains an infinite number of trunks. Thus, at no time will there be insufficient central office equipment to handle a demand for service, and no provision need be made for dealing with demands that cannot be satisfied.

The original work on this particular model for telephone traffic is in Palm,⁹ and Palm's results have been reported by Feller³ and Jensen.⁴ The results have been extended heuristically to arbitrary absolutely continuous holding-time distributions by Riordan,¹¹ following some ideas of Newland⁸ suggested by S. O. Rice.

Let $P_{ij}(t)$ be the probability that there are j trunks busy at t if there were i busy at 0. And let $P_i(t, x)$ be the generating function of these probabilities, defined by

$$P_i(t, x) = \sum_{j=0}^{\infty} x^j P_{ij}(t).$$

Then Palm⁹ has shown (pp. 56 et seq.) that

$$P_i(t, x) = [1 + (x - 1) e^{-\gamma t}]^i \exp \{ (x - 1) ah (1 - e^{-\gamma t}) \}.$$

This is formula (12) of Riordan¹¹ with his g replaced by $e^{-\gamma t}$. It can be verified that the random variable $N(t)$ is Markovian; the limit of $P_i(t, x)$ as $t \rightarrow \infty$ is

$$\exp \{ (x - 1) ah \},$$

so that the equilibrium distribution of the number of trunks in use is a Poisson distribution with mean $b = ah$. The shifted random variable $[N(t) - b]$ then has mean zero, and covariance function $be^{-\gamma t}$.

For additional work on this model the reader is referred to F. W. Rabe,¹⁰ and to H. Stormer.¹²

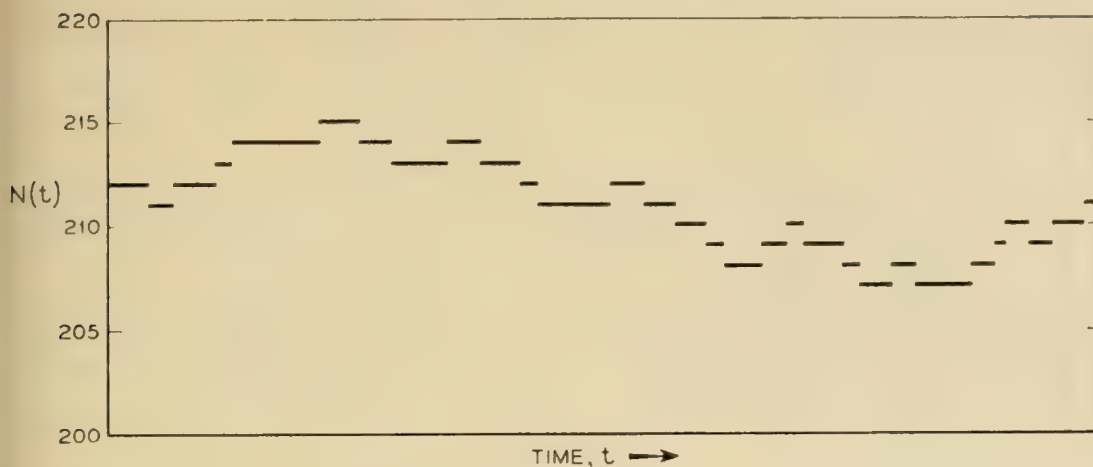
III DISCUSSION OF THE MODEL

Let us envisage the operation of the model we have described by considering the random variable $N(t)$ equal to the number of trunks busy at time t . As a random function of time, $N(t)$ jumps up one unit step each time a demand for service occurs, and it jumps down one unit step each time a conversation ends. If $N(t)$ reaches zero, it stays there until there is another demand for service. If $N(t) = n$, the probability that a conversation ends in the next small time-interval Δt is

$$n\gamma\Delta t + o(\Delta t),$$

because the n conversations are mutually independent. A graph of a sample of $N(t)$ is shown in Fig. 1.

The model we described departs from reality in several important ways, which it is well to discuss. First, the assumption that the number of trunks is infinite is not realistic, and is justified only by the mathematical complication which results when we assume the number of trunks

Fig. 1 — A graph of $N(t)$.

to be finite. It can also be argued that unlimited office capacity is approached by offices with adequate facilities and low calling rates, and therefore, in some practical cases at least, the model is not flagrantly inaccurate.

Second, the choice of a constant calling rate for the model ignores the fact that in most offices the calling rate is periodic. Thus, the applicability of our results to offices whose calling rates undergo drastic changes in time is restricted to intervals during which the normally variable calling rate is nearly constant. Finally, although the assumption of a negative exponential distribution of holding-time affords the model great mathematical convenience, it is doubtful whether in a realistic model the most likely holding-time would have length zero, as it does in the present one.

IV SUMMARY OF NOTATIONS

a = Poisson calling rate

h = mean holding-time

$\gamma = h^{-1}$ = hang-up rate per talking subscriber

$b = ah$ = average number of busy trunks

$N(t)$ = number of trunks in use at t

(O, T) = interval of observation

$n = N(O)$ = number of trunks in use at the start of observation

A = number of calls arriving in (O, T)

H = number of hang-ups in (O, T)

$K = A + H$

$$Z = \int_0^T N(t) dt$$

$$M = Z/T = \text{average of } N(t) \text{ over } (0, T)$$

$\{p_n\}$ = the (discrete) probability distribution of n , the number of trunks found busy at the start of observation

An estimator for a parameter is denoted by adding a cap (^) and a subscript. The subscripts differentiate among various estimators for the same parameter. We use $\hat{a}_c = A/T$, $\hat{\gamma}_c = H/Z$, $\hat{a}_1 = K/2T$, $\hat{\gamma}_1 = K/2Z$, and $\hat{\gamma}_2 = A/Z$.

Also, it is convenient to use the following abbreviations: r for γT , and C for $(1 - e^{-r})/r$, where r is the dimensionless ratio of observation-time to mean holding-time. The symbol E is used throughout to mean mathematical expectation.

V THE AVERAGE TRAFFIC

We have adopted a model which depends on two parameters, the calling rate a , and the mean holding-time h , or its reciprocal γ . Before searching for a set of statistics that is sufficient for the estimation of these parameters, let us consider the product $ah = b$. This product is important because, as we saw in Section II, the equilibrium distribution of the number of trunks in use depends only on b , and not on a and h individually. Indeed, the equilibrium probability that n trunks are busy is

$$\frac{e^{-b} b^n}{n!},$$

and the average number of busy trunks in equilibrium is just b .

The average number of trunks busy during a time interval T is

$$M = \frac{1}{T} \int_0^T N(t) dt;$$

i.e., the integral of the random function $N(t)$ over the interval T , divided by T . This suggests that for large time intervals T , M will come close to the value of b , and can be used as an estimator of b . Since M is a random variable, the question arises, what are the statistical properties of M ? This question has been considered in the literature, the principal references being to F. W. Rabe¹⁰ and to J. Riordan.¹¹ Riordan's paper is a determination of the first four semi-invariants of the distribution of M during a period of statistical equilibrium, but without restriction on the

assumed frequency distribution of holding-time. It follows from Rior-dan's results that M converges to b in the mean, which is to say that

$$\lim_{T \rightarrow \infty} E \{ |M - b|^2 \} = 0.$$

It also follows that M is an unbiased estimator of b ; i.e., that $E\{M\} = b$, and that M is a consistent estimator of b , which means that

$$\lim_{T \rightarrow \infty} pr \{ |M - b| > \varepsilon \} = 0$$

for each $\varepsilon > 0$.

VI MAXIMUM CONDITIONAL LIKELIHOOD ESTIMATORS

As shown in Appendix A, the likelihood L_c of an observed sequence, conditional on $N(O)$, is defined by

$$\ln L_c = A \ln a + H \ln \gamma - \gamma Z - aT.$$

According to the method of maximum likelihood, we should select, as estimators of a and γ respectively, quantities $\hat{a}_c$ and $\hat{\gamma}_c$ which maximize the likelihood L_c . Now a maximum of L_c is also one of $\ln L_c$, and vice versa. Therefore a_c and γ_c are determined as roots of the following two equations, called the likelihood equations:

$$\frac{\partial}{\partial a} \ln L_c = 0; \quad \frac{\partial}{\partial \gamma} \ln L_c = 0.$$

The solutions to the likelihood equations are

$$\hat{a}_c = \frac{A}{T}, \quad \hat{\gamma}_c = \frac{H}{Z}.$$

These are the maximum conditional likelihood estimators of a and γ . The estimator $\hat{a}_c$ is the number of requests for service in T divided by T ; this is intuitively satisfactory, since $\hat{a}_c$ estimates a calling rate.

Since maximum likelihood estimators of functions of parameters are generally the same functions of maximum likelihood estimators of the parameters, we see that AZ/HT is a maximum likelihood estimator of b .

VII PRACTICAL ESTIMATORS SUGGESTED BY MAXIMIZING THE LIKELIHOOD L , DEFINED IN APPENDIX A

We obtain as likelihood equations

$$\frac{\partial}{\partial a} \ln L = 0, \quad \frac{\partial}{\partial \gamma} \ln L = 0.$$

These may be written as

$$a = \frac{n + A}{h + T},$$

and

$$\gamma = \frac{H + \frac{a}{\gamma}}{Z + \frac{n}{\gamma}}.$$

The first of these shows the estimated calling rate as a pooled combination of the conditional estimate A/T , considered in the last section, and an estimate n/h based on the initial state. This latter estimate has the form

$$\frac{\text{calls in progress}}{\text{mean holding time}},$$

and so is intuitively reasonable, since $b/h = a$. The second equation exhibits our estimate of γ as a pooled combination of the conditional estimate H/Z and the ratio a/n . This ratio is acceptable as an estimate of γ , since $a/b = \gamma$ and $b = E\{n\}$ is the average value of n .

If we substitute, in the right-hand sides of these equations, the conditional estimators A/T , H/Z , and Z/H for a , γ , and h , respectively, we obtain simple, intuitive estimators which include the influence of the initial state n , and show how it decreases with increasing T . Thus

$$\frac{n + A}{\frac{Z}{H} + T} \quad \text{estimates } a,$$

$$\frac{H + \frac{AZ}{TH}}{Z + \frac{nZ}{H}} \quad \text{estimates } \gamma.$$

VIII OTHER ESTIMATORS

Additional estimators may be arrived at by intuitive considerations, or by modifying certain maximum likelihood estimators. Some estimators so obtained are important because they use more of the information available in an observation than do the conditional estimators $\hat{a}_c$ and $\hat{\gamma}_c$, without being so complicated functionally that we cannot easily study their statistical properties.

It seems reasonable, and can be shown rigorously (Appendix C), that for an interval (O, T) of statistical equilibrium, the distribution of A and that of H are the same. Thus we can argue that, for long time intervals, A and H will not be very different. This suggests using

$$\hat{a}_1 = \frac{A + H}{2T} = \frac{K}{2T}$$

as an estimator of a . This estimator does not involve γ , and it uses not only information given by A , but also information supplied by arrivals occurring possibly before the start of observation.

Similarly, since $b = a/\gamma$, and M is a consistent and unbiased estimator of b , we may use

$$\hat{\gamma}_1 = \frac{K}{2Z} = \frac{1}{\hat{h}_1}$$

to estimate γ , and its reciprocal to estimate h . Finally, since for long (O, T) we have $A \sim H$, we may try

$$\hat{\gamma}_2 = \frac{A}{Z} = \frac{1}{\hat{h}_2}$$

as an estimator of γ , and its reciprocal as an estimator of h .

IX THE JOINT DISTRIBUTION OF THE SUFFICIENT STATISTICS

The basic result of this paper is a formula for the generating function

$$E\{z^n x^{N(T)} w^A u^H e^{-\zeta Z}\} \quad (9.1)$$

for the joint distribution of the random variables n , $N(T)$, A , H , and Z . This formula is derived in Appendix C, by methods illustrated in Section X. For an initial n distribution $\{p_n\}$, the generating function is

$$\sum_{n \geq 0} p_n z^n \left[\frac{(\zeta x + \gamma x - \gamma u) e^{-(\zeta + \gamma)T} + \gamma u}{\zeta + \gamma} \right]^n \cdot \exp \left\{ \frac{aw(\zeta x + \gamma x - \gamma u)[1 - e^{-(\zeta + \gamma)T}]}{(\zeta + \gamma)^2} + \frac{a\gamma w u T}{\zeta + \gamma} - aT \right\}. \quad (9.2)$$

It is proved in Appendix A that the set of statistics $\{n, A, H, Z\}$ is sufficient for estimation on the basis of the information assumed, which was described in Section I. Thus the generating function (9.2) specifies, at least in principle, what can be discovered about the process from an observation interval (O, T) , for which $N(O)$ has the distribution $\{p_n\}$. All the results summarized in this section are consequences of (9.2).

TABLE II

X	$\ln E\{X\}$
1. $e^{-\zeta Z}$	$b \left[-\zeta T + \frac{\zeta^2 T}{\zeta + \gamma} - \frac{\zeta^2 (1 - e^{-(\zeta + \gamma) T})}{(\zeta + \gamma)^2} \right]$
2. $e^{-\zeta M}$	$b \left[-\zeta + \frac{\zeta^2}{\zeta + r} - \frac{\zeta^2 (1 - e^{-(\zeta + r)})}{(\zeta + r)^2} \right]$
3. y^K	$2aTC(y - 1) + aT(1 - C)(y^2 - 1)$
4. $e^{-\zeta \hat{a}_1}$	$2aTC(e^{-\zeta/2 T} - 1) + aT(1 - C)(e^{-\zeta/T} - 1)$
5. $y^K e^{-\zeta M}$	$b \left[\left(1 - \frac{ry}{\zeta + r} \right)^2 [e^{-\zeta (+r)} - 1] - r \left(1 - \frac{ry^2}{\zeta + r} \right) \right]$

By substitution, and by either letting the appropriate power series variables $\rightarrow 1$, or letting $\zeta \rightarrow 0$, or both, we can obtain from (9.2) the generating function of any combination of linear functions of the basic random variables n , $N(T)$, A , H , and Z . Some of the generating functions thereby obtained are listed in Table II, in which the entries all refer to an interval (O, T) of equilibrium.

Since, for equilibrium (O, T) , the generating functions are all exponentials, it has been convenient to make Table II a table of logarithms of expectations, with random variables X on the left, and functions $\ln E\{X\}$ on the right. C as a function of r is plotted in Fig. 2.

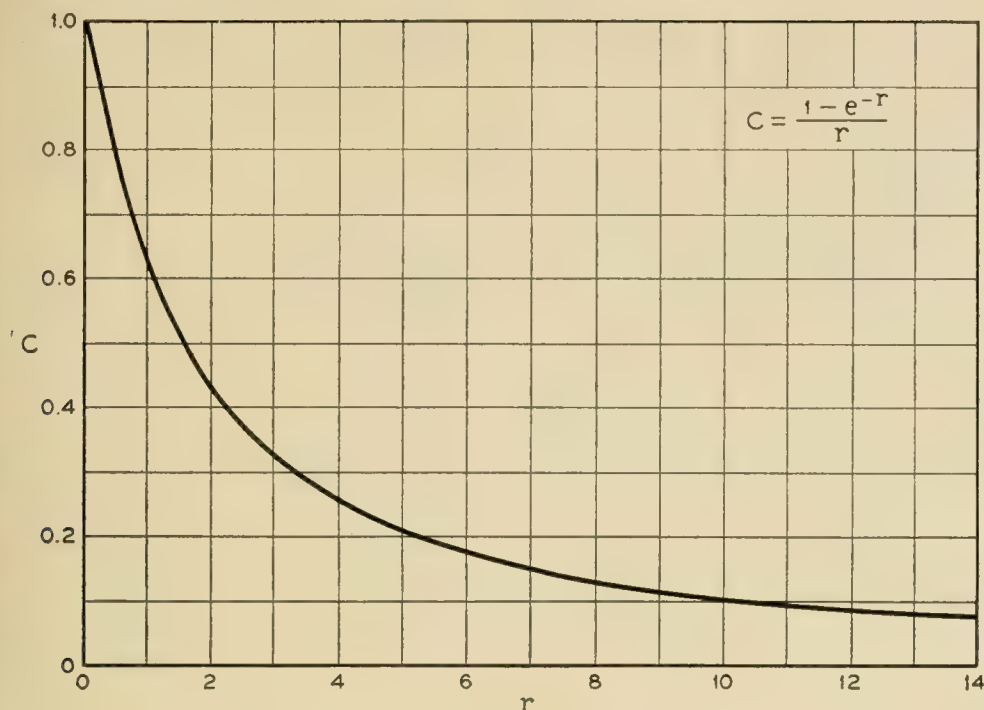
Entry 1 of Table II is actually the cumulant generating function of Z for equilibrium (O, T) ; similarly, Entry 2 is that of M , and depends only on the average traffic b and the ratio r . The form of the general cumulant of M is

$$k_n = b \frac{n(n - 1)}{T^n} \int_0^T (T - x)x^{n-2}e^{-\gamma x} \, dx.$$

This result coincides with a special case (exponential holding-time) of a conjecture of Riordan.¹¹ This conjecture was first established (for a general holding-time distribution) in unpublished work of S. P. Lloyd. The cumulant generating function permits investigation of asymptotic properties. We prove in Section X that the standardized variable

$$\begin{aligned} v &= (\gamma T/2b)^{1/2} (M - b) \\ &= (r/2b)^{1/2} (M - b) \end{aligned}$$

is asymptotically normally distributed with mean 0 and variance 1.

Fig. 2 — C as a function of r .

From Entry 3 of Table II it can be seen that K is distributed as $2u + v$, where u and v follow independent Poisson distributions with the respective parameters $aT(1 - C)$ and $2aTC$. The probability that $K = n$ for an interval of equilibrium is

$$r_n = \exp \{aT(C - 1)\} \sum \frac{(2aTC)^{n-2j}}{(n - 2j)!} \frac{(aT - aTC)^j}{j!},$$

where the sum is over j 's for which $0 \leq 2j \leq n$.

The estimator $\hat{a}_1$ for a is equal to $K/2T$, and has mean and variance given by

$$E\{\hat{a}_1\} = a,$$

$$\text{var} \{\hat{a}_1\} = \frac{a}{2T} (2 - C).$$

The distribution of $\hat{a}_1$ is given by

$$\text{pr}\{\hat{a}_1 \leq x\} = \sum r_n,$$

the summation being over $n \leq 2Tx$.

From (9.2) one can obtain, by substitution of the stationary n distribution for $\{p_n\}$, and subsequent differentiation, the means, variances, covariances, and correlation coefficients of the sufficient statistics, for

TABLE III — $E\{X, Y\}$

	l	n	A	H	K	Z
l	1	b	aT	aT	$2aT$	bT
n		$b(1 + b)$	baT	$aT(C + b)$	$aT(C + 2b)$	$bT(C + b)$
A			$aT(1 + aT)$	$aT(1 - C + aT)$	$aT(2 - C + aT)$	$bT(1 - C + aT)$
H				$aT(1 + aT)$	$aT(2 - C + aT)$	$bT(1 - C + aT)$
K					$2aT(2 - C + 2aT)$	$2bT(1 - C + aT)$
Z						$bTh(2 - 2C + aT)$

TABLE IV — $\text{cov}\{X, Y\}$

	n	A	H	K	Z
n	b	0	aTC	aTC	bTC
A		aT	$aT(1 - C)$	$aT(2 - C)$	$bT(1 - C)$
H			aT	$aT(2 - C)$	$bT(1 - C)$
K				$2aT(2 - C)$	$2bT(1 - C)$
Z					$2bTh(1 - C)$

equilibrium intervals (O, T) . It has been convenient to display these in three triangular arrays, the first consisting of expectations of products, the second comprising the variances and covariances, and the third exhibiting, for simplicity, the squared correlation coefficients, since the correlation coefficients are never negative for these random variables.

In Table III, the entry with coordinates (X, Y) is $E\{XY\}$ for equilibrium (O, T) . All three tables are expressed in terms of a, b, T, h, r , and C , the last of which is plotted in Fig. 2.

The variances and covariances of the sufficient statistics are listed in Table IV; the entries are of the form:

$$\text{cov}\{X, Y\} = E\{XY\} - E\{X\}E\{Y\}.$$

Table V, finally, lists the squared correlation coefficients; i.e., the quantities

$$\rho^2(X, Y) = \frac{\text{cov}^2\{X, Y\}}{\text{var}\{X\} \text{var}\{Y\}}.$$

For any time interval (O, T) , A has a Poisson distribution with parameter aT , so that $T\hat{a}_c$ does also. Therefore the distribution of $\hat{a}_c$ is given by

$$\text{pr}\{\hat{a}_c \leq x\} = \sum \frac{e^{-aT}(aT)^n}{n!},$$

where the summation is over $n \leq xT$. Evidently

$$E\{\hat{a}_c\} = a,$$

and

$$\text{var } \{\hat{a}_c\} = \frac{a}{T},$$

so that $\hat{a}_c$ is an unbiased and consistent estimator of a . We now compare the variances of estimators $\hat{a}_c$ and $\hat{a}_1$. From Table IV we have

$$\text{var } \{\hat{a}_1\} = \frac{a}{T} \left(1 - \frac{C}{2}\right) < \frac{a}{T} = \text{var } \{\hat{a}_c\},$$

so that $\hat{a}_1$ is a better estimator of a for any $T > 0$, in the sense that its variance is less.

X THE DISTRIBUTIONS of Z AND M

Since we have defined

$$Z = \int_0^T N(t) \, dt,$$

we can regard Z as the result of growth whose rate is given by the random step-function $N(t)$; when $N(t) = n$, Z is growing at rate n . An idea similar to this is used by Kosten, Manning, and Garwood⁶, and by Kosten alone.⁵ Now the $Z(T)$ process by itself is not Markovian, but it can be seen that the two-dimensional variable $\{N(t), Z(t)\}$ itself is Markovian. Let $F_n(z, t)$ be the probability that $N(t) = n$ and $Z(t) \leq z$. Since the two-dimensional process is Markovian, we can derive infinitesimal relations for $F_n(z, t)$ by considering the possible changes in the system during a small interval of time Δt .

TABLE V — $\rho^2(X, Y)$

	n	A	H	K	Z
n	1	0	$1 - e^{-r}$	$\frac{rC^2}{2 - C}$	$\frac{rC^2}{2(1 - C)}$
A		1	$1 - C$	$\frac{2 - C}{2}$	$\frac{1 - C}{2}$
H			1	$\frac{2 - C}{2}$	$\frac{1 - C}{2}$
K				1	$\frac{1 - C}{2 - C}$
Z					1

If $N(t) = n$, then the probability is $[1 - \gamma n \Delta t - a \Delta t - o(\Delta t)]$ that there is neither a request for service nor a hang-up during Δt following t , and that $Z(t + \Delta t) = Z(t) + n \Delta t$. Therefore the conditional probability that $N(t + \Delta t) = n$ and $Z(t + \Delta t) \leq z$, given that no changes occurred in Δt , is

$$F_n(z - n \Delta t, t).$$

For $N(t) = (n + 1)$, the probability is $\gamma(n + 1) \Delta t + o(\Delta t)$ that one conversation will end during Δt following t . The increment to $Z(t)$ during Δt will depend on the length x of the interval from t to the point within Δt at which the conversation ended. The increment has magnitude $(n + 1)x + n(\Delta t - x) = x + n \Delta t$, as can be verified from Fig. 3, in which the shaded area is the increment. Since x is distributed uniformly between 0 and Δt , the increment $x + n \Delta t$ is distributed uniformly between $n \Delta t$ and $(n + 1) \Delta t$. Therefore the conditional probability that $N(t + \Delta t) = n$ and $Z(t + \Delta t) \leq z$, given that one conversation ended in Δt , is

$$\frac{1}{\Delta t} \int_{n \Delta t}^{(n+1) \Delta t} F_{n+1}(z - u, t) du.$$

By a similar argument it can be shown that the probability that one request for service arrives in Δt is $a \Delta t + o(\Delta t)$, and that the conditional probability that $N(t + \Delta t) = n$ and $Z(t + \Delta t) \leq z$, given that one request arrived during Δt , is

$$\frac{1}{\Delta t} \int_{(n-1) \Delta t}^{n \Delta t} F_{n-1}(z - u, t) du.$$

Define $F_n(z, t)$ to be identically 0 for negative n . Adding up the probabil-

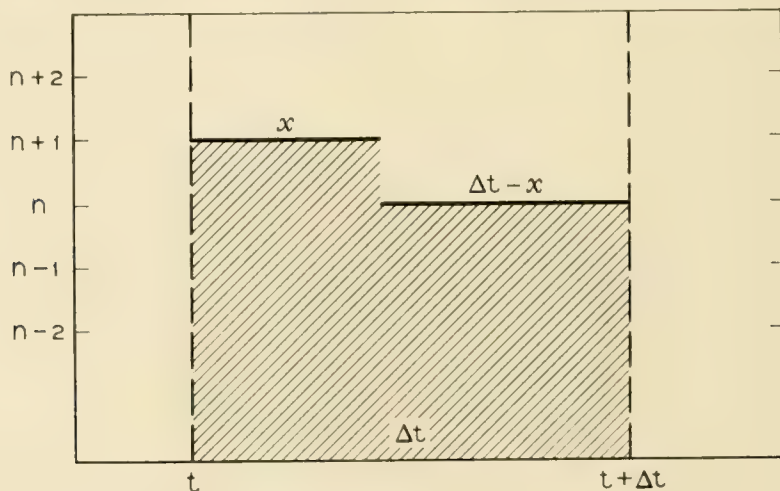


Fig. 3 — Increment to Z in Δt .

ities of mutually exclusive events, we obtain the following infinitesimal relations for $F_n(z, t)$:

$$\begin{aligned} F_n(z, t + \Delta t) = & \gamma(n + 1) \int_{n\Delta t}^{(n+1)\Delta t} F_{n+1}(z - u, t) du \\ & + a \int_{(n-1)\Delta t}^{n\Delta t} F_{n-1}(z - u, t) du + F_n(z - n\Delta t, t) \\ & \cdot [1 - \Delta t(\gamma n + a)] + o(\Delta t), \quad \text{for any } n. \end{aligned}$$

Expanding the penultimate term of the right side in powers of $n\Delta t$, and the left side in powers of Δt , we divide by Δt , and take the limit as Δt approaches 0. Now

$$\lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \int_{n\Delta t}^{(n+1)\Delta t} F_{n+1}(z - u, t) du = F_{n+1}(z, t).$$

Thus, omitting functional dependence on z and t for convenience, we reach the following partial differential equations for $F_n(z, t)$:

$$\begin{aligned} \frac{\partial}{\partial t} F_n = & \gamma(n + 1)F_{n+1} + aF_{n-1} - n \frac{\partial}{\partial z} F_n \\ & - [\gamma n + a]F_n, \quad \text{for any } n. \end{aligned} \tag{10.1}$$

Since $Z(0) = 0$, we impose the following boundary conditions:

$$\begin{aligned} F_n(0, t) &= 0 \quad \text{for } n > 0 \text{ and } t > 0, \\ F_n(z, 0) &= p_n \quad \text{for } z \geq 0, \\ F_n(z, 0) &= 0 \quad \text{for } z < 0, \end{aligned} \tag{10.2}$$

where the sequence $\{p_n\}$ forms an arbitrary $N(0)$ distribution that is zero for negative n .

To transform the equations, we introduce the Laplace-Stieltjes integrals

$$\varphi_n(\zeta, t) = \int_{0-}^{\infty} e^{-\zeta z} dF_n(z, t), \quad t \geq 0, \quad \operatorname{Re}(\zeta) > 0,$$

in which the Stieltjes integration is understood always to be on the variable z . We note that

$$\int_{0-}^{\infty} e^{-\zeta z} F_n(z, t) dz = \frac{1}{\zeta} \varphi_n(\zeta, t),$$

and that

$$\varphi_n(\zeta, t) = F_n(0, t) + \int_0^{\infty} e^{-\zeta z} \frac{\partial}{\partial z} F_n(z, t) dz.$$

Applying now the Laplace-Stieltjes transformation to (10.1), we obtain

$$\frac{\partial \varphi_n}{\partial t} = \gamma(n+1)\varphi_{n+1} + a\varphi_{n-1} - n\zeta\varphi_n + n\zeta F_n(0, t) - [\gamma n + a]\varphi_n, \quad (10.3)$$

in which we have left out functional dependence on ζ and t where it is unnecessary. By the boundary conditions (10.2), $n\zeta F_n(0, t) = 0$ for $n \geq 0$ and $t > 0$; in (10.3) we may therefore omit this term in the region $t > 0$. Let φ be defined by

$$\varphi(x, \zeta, t) = \sum_{n=0}^{\infty} x^n \varphi_n(\zeta, t).$$

The series is absolutely convergent for $|x| < 1$, since

$$|\varphi_n(\zeta, t)| \leq 1, \text{ for all } n.$$

The following partial differential equation for φ is obtained from (10.3):

$$\frac{\partial \varphi}{\partial t} + [\zeta x + \gamma(x-1)] \frac{\partial \varphi}{\partial x} = a(x-1)\varphi. \quad (10.4)$$

If we integrate out the information about Z by letting ξ approach 0 in this equation, we obtain the equation derived by Palm (loc. cit.) for the generating function of $N(t)$. Therefore our equation has a solution of the same form as Palm's. For the boundary conditions (10.2), this solution is

$$\varphi = \exp \left\{ \frac{a[1 - e^{-(\zeta+\gamma)t}]}{(\zeta + \gamma)^2} [\zeta x + \gamma(x-1)] - \frac{a\zeta t}{\zeta + \gamma} \right\} \cdot \sum_{n=0}^{\infty} p_n \left[\frac{[\zeta x + \gamma(x-1)]e^{-(\zeta+\gamma)t} + \gamma}{\zeta + \gamma} \right]^n. \quad (10.5)$$

Actually φ contains more information than we want since it yields the joint distribution of N and Z . We may integrate out the former variable by letting x approach 1 in 10.5. Then,

$$E\{\exp(-\zeta Z)\} = \exp \left\{ \frac{a\zeta(1 - e^{-(\zeta+\gamma)T})}{(\zeta + \gamma)^2} - \frac{a\zeta T}{\zeta + \gamma} \right\} \cdot \sum_{n=0}^{\infty} p_n \left[\frac{\zeta e^{-(\zeta+\gamma)T} + \gamma}{\zeta + \gamma} \right]^n$$

is the Laplace transform of the distribution of Z for an arbitrary $N(O)$ distribution $\{p_n\}$. This result is not restricted to an interval (O, T)

of statistical equilibrium; however, if the sequence $\{p_n\}$ does form the stationary distribution discussed in Section II, then

$$\sum x^n p_n = \exp \{b(x - 1)\}, \quad (10.6)$$

and

$$\psi = \exp \left\{ \frac{b\xi^2(e^{-(\xi+\gamma)T} - 1)}{(\xi + \gamma)^2} - \frac{a\xi T}{\xi - \gamma} \right\} \quad (10.7)$$

is the Laplace transform of the distribution of Z for an interval $(0, T)$ of statistical equilibrium.

The Laplace transform is a moment generating function expressible as

$$\psi = \sum_{n=0}^{\infty} \frac{(-\xi)^n m_n}{n!},$$

where m_n is the n^{th} ordinary moment of Z . Differentiation of 10.7 then gives a recurrence relation for the moments upon equating powers of $(-\xi)$. Thus,

$$\begin{aligned} (\xi + \gamma)^3 \left(-\frac{\partial \psi}{\partial \xi} \right) \\ = \psi \cdot \{2a\xi(1 - e^{-(\xi+\gamma)T}) + (\xi + \gamma)[\gamma aT + bT\xi^2 e^{-(\xi+\gamma)T}]\}, \end{aligned}$$

and

$$\begin{aligned} \gamma^3 m_{n+1} - 3\gamma^2 n m_n + 3\gamma n(n-1)m_{n-1} - n(n-1)(n-2)m_{n-2} \\ = a\gamma^2 T m_n - (2a + a\gamma T)nm_{n-1} + 2ane^{-\gamma T}(m+T)^{n-1} + n \\ \cdot (n-1)aTe^{-\gamma T}(m+T)^{n-2} - n(n-1)(n-2)bTe^{-\gamma T}(m+T)^{n-3}, \end{aligned} \quad (10.8)$$

where $(m+T)^n$ is the usual symbolic abbreviation of

$$\sum_{j=0}^n \binom{n}{j} T^j m_{n-j}.$$

From the recurrence (10.8) it is easily verified that

$$\begin{aligned} m_1 &= bT, \\ m_2 &= (bT)^2 + \frac{2bT}{\gamma} [1 - C], \end{aligned}$$

from which it follows that the variance of Z is

$$\text{var } \{Z\} = \frac{2bT}{\gamma} [1 - C].$$

Since ψ is the Laplace-Stieltjes transform of the distribution of Z over an interval of equilibrium, $\ln \psi$ is the cumulant generating function, and has the following simple form:

$$\begin{aligned}\ln \psi &= b \left[\frac{\xi^2(e^{-(\xi+\gamma)T} - 1)}{(\xi + \gamma)^2} - \frac{\gamma \xi T}{\xi + \gamma} \right] \\ &= b \left[-\xi T + \frac{\xi^2 T}{\xi + \gamma} - \frac{\xi^2(1 - e^{-(\xi+\gamma)T})}{(\xi + \gamma)^2} \right].\end{aligned}\quad (10.9)$$

M is a linear function of Z , so we may obtain the cumulant generating function of M in accordance with Cramér¹ (p. 187). This function is

$$b \left[-\xi + \frac{\xi^2}{r + \xi} - \frac{\xi^2(1 - e^{-r}e^{-\xi})}{(r + \xi)^2} \right], \quad (10.10)$$

and depends only on b and r .

The mean and variance of M for an interval of equilibrium are respectively given by

$$\begin{aligned}E\{M\} &= b, \\ \text{var } \{M\} &= \frac{2b}{r} [1 - C], \quad \text{with } C = \frac{1 - e^{-r}}{r},\end{aligned}$$

results which were first proved in Riordan.¹¹ A normal distribution having the mean and variance of M has the cumulant generating function

$$b \left[-\xi + \frac{\xi^2}{r} + \frac{\xi^2(e^{-r} - 1)}{r^2} \right], \quad (10.11)$$

which is to be compared to (10.10). Since $\text{var } \{M\}$ goes to 0 as T approaches ∞ , we may expect that a suitably normalized version of Z will be asymptotically normally distributed as T approaches ∞ . The cumulant generating function of the normalized variable $(2bhT)^{-1/2}(Z - bT)$ is

$$\frac{\xi^2}{\xi \left(\frac{2}{aT} \right)^{1/2} + 2} \left[1 + \frac{\exp \left\{ -\xi \left(\frac{2b}{r} \right)^{-1/2} - r \right\} - 1}{\xi \left(\frac{2b}{r} \right)^{-1/2} + r} \right],$$

which approaches $\xi^2/2$ as $T \rightarrow \infty$. It follows that the normalized variable is asymptotically normal with mean 0 and variance 1, and that $(r/2b)^{1/2}(M - b)$ is also asymptotically normal (0, 1).

APPENDIX A

PROOF THAT $\{n, A, H, Z\}$ IS SUFFICIENT.

We observe the system during the interval (O, T) , and gather the information specified in Section I, and summarized in Table I. From this information we can extract four sets of numbers, described as follows:

S_a the set of complete observed inter-arrival times, not counting the interval from the last arrival until T

S_h the set of complete observed holding times

S_1 the set of hang-up times for calls of category (i)

S_4 the set of calling-times for calls of category (iv)

In addition, our data enable us to determine the following numbers:

n the number $N(O)$ of calls found at the start of observation

k the number of calls of category (iii); i.e., of calls which last throughout the interval (O, T)

x the length of the time-interval between the last observed arrival and T

In view of the negative exponential distributions which have been assumed for the inter-arrival times and the holding-times, and in view of the assumptions of independence, we can write the likelihood of an observed sequence of events as

$$L = e^{-k\gamma T - ax} p_n \cdot \prod_{u \in S_a} a e^{-au} \cdot \prod_{z \in S_h} \gamma e^{-\gamma z} \cdot \prod_{w \in S_1} e^{-\gamma w} \cdot \prod_{y \in S_4} e^{-\gamma(T-y)},$$

so that

$$\begin{aligned} \ln L = & -\gamma k T - ax + \ln p_n + A \ln a - \sum_{u \in S_a} au \\ & + H \ln \gamma - \sum_{z \in S_h} \gamma z - \sum_{w \in S_1} \gamma w - \sum_{y \in S_4} \gamma(T-y) \end{aligned}$$

It is easily seen that the summations and the two initial terms can be combined into a single term, so that we obtain

$$\ln L = \ln p_n + A \ln a + H \ln \gamma - \gamma Z - aT.$$

This shows that L depends only on the statistics n, A, H , and Z ; it follows that the information we have assumed can be replaced by the set of statistics $\{n, A, H, Z\}$, and that these are sufficient for estimation based on that information.

The likelihood is sometimes defined without reference to the initial state, by leaving the factor p_n out of the expression for L . Strictly speaking, this omission defines the conditional likelihood for the observed

sequence, conditional on starting at n . We use the notation:

$$L_c = \frac{L}{p_n}.$$

A definition of likelihood as L_c has been used by Moran.⁷ Clearly

$$\ln L_c = A \ln a + H \ln \gamma - \gamma Z - aT.$$

APPENDIX B

UNCONDITIONAL MAXIMUM LIKELIHOOD ESTIMATES

The definition of likelihood as L leads to complicated results which are of theoretical rather than practical interest. For this reason these results have been relegated to an appendix.

The results of setting $\partial/\partial\gamma \ln L$ and $\partial/\partial a \ln L$ equal to zero lead, respectively, to the likelihood equations

$$a - \gamma(n - H) - \gamma^2 Z = 0,$$

$$\gamma n - a + \gamma A - a\gamma T = 0.$$

Considered as a system of equations for γ and a , this pair has the non-negative roots

$$\hat{\gamma} = \frac{H - n - M + \{(H - n - M)^2 + 4MK\}^{1/2}}{2Z},$$

$$\hat{a} = \frac{K}{T} - \hat{\gamma}M.$$

These are the unconditional maximum likelihood estimators for γ and a . Although $\hat{a}_c$ depended only on A and T , and $\hat{\gamma}_c$ only on H and Z , the unconditional estimators depend on all of n , A , H , Z , and T . We may obtain a maximum unconditional likelihood estimator for b as well, either by considering L to be a function of b and γ , or from general properties of maximum likelihood estimators. Since $b = a/\gamma$, we expect that $\hat{b} = \hat{a}/\hat{\gamma}$, as can be verified by an argument similar to that used above for $\hat{a}$ and $\hat{\gamma}$.

The estimators $\hat{a}$, b , and $\hat{\gamma}$ obtained in this Appendix may turn out to be useful in practice, but their complicated dependence on the sufficient statistics n , A , H , and Z makes a study of their statistical properties difficult. As a first step along such a study, we have derived the generating function of the joint distribution of the sufficient statistics in Appendix C. The greater simplicity of the conditional estimators of Section VI makes it possible to study their statistical properties. This

fact gives them a practical ascendancy over the unconditional estimators, even though the latter may be more efficient statistically by dint of using all the information available in an observation.

APPENDIX C

THE JOINT DISTRIBUTION OF $N(t)$, n , A , H , AND Z

By methods already used in Section X one can obtain a generating function for the joint distribution of all the random variables n , $N(t)$, A , H , and Z . Let

$$\Phi = E\{x^{N(t)}w^A u^H e^{-\zeta Z}\}.$$

Then Φ satisfies the differential equation

$$\frac{\partial \Phi}{\partial t} + [\zeta x + \gamma x - \gamma u] \frac{\partial \Phi}{\partial x} = a(wx - 1)\Phi,$$

whose solution has the form

$$\Phi = R\{[\zeta x + \gamma x - \gamma u]e^{-(\zeta+\gamma)t}\} \cdot \exp\left(\frac{aw[\zeta x + \gamma x - \gamma u][1 - e^{-(\zeta+\gamma)t}]}{(\zeta + \gamma)^2} + \frac{a\gamma wut}{\zeta + \gamma} - at\right),$$

where the function R is determined by the initial distribution $\{p_n\}$ through the relation

$$R\{\xi\} = \sum_{n \geq 0} p_n \left[\frac{\xi + \gamma u}{\zeta + \gamma} \right]^n.$$

From these results it follows that the generating function

$$E\{z^n x^{N(T)} w^A u^H e^{-\zeta Z}\}$$

is given by

$$\sum_{n \geq 0} p_n z^n \left(\frac{(\zeta x + \gamma x - \gamma u)e^{-(\zeta+\gamma)T} + \gamma u}{\zeta + \gamma} \right)^n \cdot \exp\left(\frac{aw(\zeta x + \gamma x - \gamma u)[1 - e^{-(\zeta+\gamma)T}]}{(\zeta + \gamma)^2} + \frac{a\gamma wuT}{\zeta + \gamma} - aT\right).$$

If $\{p_n\}$ forms the stationary distribution, this reduces to

$$\exp \left[b \left(\frac{z(\zeta x + \gamma x - \gamma u)e^{-(\zeta+\gamma)T} + \gamma uz}{\zeta + \gamma} - 1 \right) + \frac{aw(\zeta x + \gamma x - \gamma u)[1 - e^{-(\zeta+\gamma)T}]}{(\zeta + \gamma)^2} + \frac{a\gamma wuT}{\zeta + \gamma} - aT \right].$$

If, in this last expression, we let x approach 1, z approach 1, and u approach 1, we obtain

$$\exp \left[\left(1 - \frac{\gamma w}{\zeta + \gamma} \right) \left(\frac{b\zeta(e^{-(\zeta+\gamma)T} - 1)}{\zeta + \gamma} - aT \right) \right] \quad (C)$$

as the generating function $E\{w^A e^{-\zeta Z}\}$ for an interval of equilibrium. Alternately, if instead we let x approach 1, z approach 1, and w approach 1, we obtain (C) with u substituted for w ; this implies the not-surprising result that for an interval of equilibrium, the two-dimensional variables $\{A, Z\}$ and $\{H, Z\}$ have the same distribution. From this and (C) it follows that for equilibrium (O, T) , (i) A and H both have a Poisson distribution with mean aT , and (ii) the estimators $\hat{h}_c$ and $\hat{h}_2$ have the same distribution.

ACKNOWLEDGEMENT

The author would like to express his gratitude for the helpful comments and suggestions of E. N. Gilbert, S. P. Lloyd, J. Riordan, J. Tukey, and P. J. Burke.

REFERENCES

1. H. Cramér, *Mathematical Methods of Statistics*, Princeton, 1946.
2. W. Feller, *An Introduction to Probability Theory and its Applications*, John Wiley and Sons, New York, 1950.
3. W. Feller, On the Theory of Stochastic Processes with Particular Reference to Applications, *Proceedings of the Berkeley Symposium on Math. Statistics and Probability*, Univ. of California Press, 1949.
4. A. Jensen, An Elucidation of Erlang's Statistical Works Through the Theory of Stochastic Processes, in *The Life and Works of A. K. Erlang*, Copenhagen, 1948.
5. L. Kosten, On the Accuracy of Measurements of Probabilities of Delay and of Expected Times of Delay in Telecommunication Systems, *App. Sci. Res.*, **B2**, pp. 108-130 and pp. 401-415, 1952.
6. L. Kosten, J. R. Manning, F. Garwood, On the Accuracy of Measurements of Probabilities of Loss in Telephone Systems, *Journal of the Royal Statistical Society (B)*, **11**, pp. 54-67, 1949.
7. P. A. P. Moran, Estimation Methods for Evolutive Processes, *Journal of the Royal Statistical Society (B)*, **13**, pp. 141-146, 1951.
8. W. F. Newland, A Method of Approach and Solution to Some Fundamental Traffic Problems, *P.O.E.E. Journal*, **25**, pp. 119-131, 1932-1933.
9. C. Palm, Intensitätsschwankungen im Fernspreckverkehr, *Ericsson Technics*, **44**, 1943.
10. F. W. Rabe, Variations of Telephone Traffic, *Elec. Comm.*, **26**, 243-248, 1949.
11. J. Riordan, Telephone Traffic Time Averages, *B.S.T.J.*, **30**, 1129-1144, 1951.
12. H. Stormer, Anwendung des Stichprobenverfahrens beim Beurteilen von Fernspreckverkehrsmessungen, *Archiv der Elektrischen Übertragung*, **8**, pp. 439-436, 1954.

Fluctuations of Telephone Traffic

By V. E. BENEŠ

(Manuscript received November 9, 1956)

The number of calls in progress in a simple telephone exchange model characterized by unlimited call capacity, a general probability density of holding-time, and randomly arriving calls is defined as $N(t)$. A formula, due to Riordan, for the generating function of the transition probabilities of $N(t)$ is proved. From this generating function, expressions for the covariance function of $N(t)$ and for the spectral density of $N(t)$ are determined. It is noted that the distributions of $N(t)$ are completely specified by the covariance function.

I INTRODUCTION

The aim of this paper is to study the average fluctuations of telephone traffic in an exchange, by means of a simple mathematical model to which we apply concepts used in the theory of stochastic processes and in the analysis of noise.

The mathematical model we use is based on the following assumptions: (1) requests for telephone service arise individually and collectively at random at an average rate of a per second; (2) the holding-times of calls are mutually independent random variables having the common probability density function $h(u)$; and (3) the capacity of the exchange is effectively unlimited, and no call is blocked or delayed by lack of equipment. This telephone exchange model has been described by J. Riordan.⁵

As a measure of traffic, it is natural to use the number of calls in progress in the exchange. We are thus led to consider a random step-function of time $N(t)$, defined as the number of calls in progress at time t . $N(t)$ fluctuates about an average in a manner depending on the calling-rate, a , and the holding-time density, $h(u)$.

II PROOF OF RIORDAN'S FORMULA FOR TRANSITION PROBABILITIES

Let $P_{m,n}(t)$ be the probability that n calls are in progress at t if m calls were in progress at 0. Define the generating function of these prob-

abilities as

$$P_m(t, x) = \sum_{n \geq 0} P_{m,n}(t) x^n,$$

and let

$$f(u) = \int_u^\infty h(x) dx,$$

so that the average holding-time, h , is given by

$$h = \int_0^\infty f(u) du.$$

Riordan⁵ has given the following formula for $P_m(t, x)$:

$$P_m(t, x) = [1 + (x - 1)g(t)]^m \exp \{ (x - 1)ah[1 - g(t)] \}, \quad (1)$$

with

$$g(t) = \frac{1}{h} \int_t^\infty f(u) du.$$

For exponential holding-time density, this formula had already been derived (as the solution of a differential equation) by Palm.²

In private communication, J. Riordan has suggested that his proof of (1) is incomplete. We therefore give a new proof of (1).

We seek the generating function of $N(t)$, conditional on the event $N(0) = m$. We obtain it by first computing the joint generating function of $N(0)$ and $N(t)$; that is,

$$E\{y^{N(0)}x^{N(t)}\}. \quad (2)$$

The desired conditional generating function is then the coefficient of y^m in (2), divided by the probability that $N(0) = m$.

To obtain a formula for (2), we exhaust the interval $(-\infty, 0)$ by division into a countable set of disjoint intervals, I_n , the n^{th} having length $T_n > 0$. Let S_n be the sum of the first n lengths, T_j . Let $\xi_n(t)$, for $t > -S_{n-1}$, be the number of those calls which arrive in I_n and are still in progress at t . And let $\eta(t)$ be the number of calls arriving during $(0, t)$, $t > 0$, and still in existence at t . Then

$$N(0) = \sum_{n \geq 1} \xi_n(0), \quad (3)$$

$$N(t) = \eta(t) + \sum_{n \geq 1} \xi_n(t), \quad t > 0. \quad (4)$$

Since calls arriving during disjoint intervals are independent, we know

that $\eta(t)$ is independent of all the ξ 's, and that $\xi_n(t)$ is independent of $\xi_j(\tau)$ if $n \neq j$. Of course, $\xi_n(t)$ and $\xi_n(\tau)$ are not independent. It follows that if the infinite product converges, then for $t > 0$

$$E\{y^{N(0)}x^{N(t)}\} = E\{x^{\eta(t)}\} \prod_{n=1}^{\infty} E\{y^{\xi_n(0)}x^{\xi_n(t)}\}. \quad (5)$$

We now compute the terms of the product. If a call originates in interval I_n , it still exists at 0 with probability

$$Q_n = \frac{1}{T_n} \int_0^{T_n} f(u + S_{n-1}) du = \frac{1}{T_n} \int_{S_{n-1}}^{S_n} f(u) du.$$

Hence if k calls arrived in I_n , the probability that m of them are still in progress at 0 is

$$\begin{aligned} \text{pr}\{\xi_n(0) = m \mid k \text{ calls arrive in } I_n\} \\ = \binom{k}{m} Q_n^m (1 - Q_n)^{k-m}, \quad m \leq k. \end{aligned}$$

Similarly, if a call originates in I_n and exists at 0, it also exists at $t > 0$ with probability

$$K_n = (Q_n T_n)^{-1} \int_0^{T_n} f(u + t + S_{n-1}) du.$$

Therefore

$$\begin{aligned} E\{x^{\xi_n(t)} \mid \xi_n(0) = m \text{ and } k \text{ calls arrive in } I_n\} \\ = [1 + (x - 1)K_n]^m, \end{aligned}$$

and so

$$\begin{aligned} E\{y^{\xi_n(0)}x^{\xi_n(t)} \mid k \text{ calls arrive in } I_n\} \\ = \{1 + \langle y[1 + (x - 1)K_n] - 1 \rangle Q_n\}^k \\ = \alpha^k. \end{aligned}$$

The number of calls arriving during I_n has a Poisson distribution with mean aT_n ; hence

$$\begin{aligned} E\{y^{\xi_n(0)}x^{\xi_n(t)}\} &= \exp\{aT_n(\alpha - 1)\} \\ &= \exp\{aT_n Q_n \langle y[1 + (x - 1)K_n] - 1 \rangle\}. \end{aligned} \quad (6)$$

By reasoning like that leading to (6), it can be shown that

$$\begin{aligned}
 E\{x^{n(t)}\} &= \exp \left\{ at(x-1) \frac{1}{t} \int_0^t f(u) du \right\} \\
 &= \exp \{ ah(x-1)[1-g(t)] \}.
 \end{aligned} \tag{7}$$

Now

$$\begin{aligned}
 \sum_{n \geq 1} aT_n Q_n &= a \int_0^\infty f(u) du = ah, \\
 \sum_{n \geq 1} aT_n Q_n K_n &= a \sum_{n \geq 1} \int_0^{T_n} f(u+t+S_{n-1}) du, \\
 &= a \int_t^\infty f(u) du = ahg(t).
 \end{aligned}$$

Therefore the infinite product is convergent, and

$$\begin{aligned}
 E\{y^{N(0)} x^{N(t)}\} &= \exp \{ ah(x-1)[1-g(t)] \\
 &\quad + \sum_{n \geq 1} aT_n Q_n \langle y[1+(x-1)K_n] - 1 \rangle \} \\
 &= \exp \{ ah \langle (x-1)[1-g(t)] + (y-1) + y(x-1)g(t) \rangle \}.
 \end{aligned} \tag{8}$$

Thus the generating function of the joint distribution of $N(0)$ and $N(t)$ is independent of the division of $(-\infty, 0)$ into intervals I_n . By letting x approach 1 in (8) and finding the coefficient of y^m in the resulting limit, we find that

$$\text{pr}\{N(0) = m\} = \frac{e^{-ah}(ah)^m}{m!}. \tag{9}$$

The coefficient of y^m in (8) itself is

$$\frac{e^{-ah}(ah)^m}{m!} [1+(x-1)g(t)]^m \exp \{ (x-1)ah[1-g(t)] \},$$

and so using (9) we find that the required conditional generating function of $N(t)$, given $N(0) = m$, is given by Riordan's formula (1).

III THE AUTOCORRELATION

In terms of $N(t)$ one can define various stochastic integrals which will be characteristic of the process. A simple one which has been extensively treated in connection with estimating the average traffic is

$$M = \frac{1}{T} \int_0^T N(t) dt,$$

the average of $N(t)$ over an interval $(0, T)$. The chief references in the literature on M are References 3 and 5. If we consider $N(t)$ during an interval $(0, T + \tau)$, a measure of the coherence of $N(t)$ during this interval, i.e., of the extent to which $N(t)$ hangs together, is given by the integral

$$U(T, \tau) = \frac{1}{T} \int_0^T N(t) N(t + \tau) dt,$$

depending on values of $N(t)$ taken τ apart. When the limit $\psi(\tau)$ of u as T approaches ∞ exists, it is usually called the autocorrelation function; most statisticians, however, reserve the term "correlation" for suitably normalized, dimensionless quantities. It can be shown that this limit exists and is the same for almost all $N(t)$ in the ensemble. It then coincides with the ensemble average, i.e.,

$$\begin{aligned} \psi(\tau) &= \lim_{T \rightarrow \infty} U(T, \tau), \text{ almost all } N(t), \\ &= E\{N(t)N(t + \tau)\}. \end{aligned}$$

The function, ψ , for the system we are discussing is derived by Riordan,⁵ and we reproduce his argument for ease of understanding. For equilibrium, and $b = ah$, we have

$$E\{N(t)N(t + \tau)\} = \sum_{m=0}^{\infty} \frac{e^{-b}(b)^m}{m!} m \frac{\partial}{\partial x} P_m(\tau, x) \Big]_{x=1}.$$

Now

$$\frac{\partial}{\partial x} P_m(\tau, x) \Big]_{x=1} = mg(\tau) + b[1 - g(\tau)],$$

so that

$$\begin{aligned} \psi(\tau) &= \sum_{m=0}^{\infty} \frac{e^{-b}b^m}{m!} m\{mg(\tau) + b[1 - g(\tau)]\}, \\ &= b^2 + bg(\tau). \end{aligned} \tag{10}$$

(Cf.,⁵ p. 1136)

The limiting value of $\psi(\tau)$ for τ approaching ∞ is the square of the mean occupancy, b , and the limiting value of $\psi(\tau)$ for τ approaching 0 is the mean square occupancy, $b^2 + b$, the second moment of the Poisson distribution with mean b .

IV THE COVARIANCE AND SPECTRAL DENSITY

The average value of $N(t)$ is $b = ah$. One way to study the fluctuations of $N(t)$ about its average is by means of the power spectrum used in the analysis of noise. (Cf. Rice.⁴) We resolve the difference $[N(t) - b]$ into sinusoidal components of non-negative frequency, and postulate a noise current proportional to this difference dissipating power through a unit resistance. The spectrum $w(f)$ is then the average power due to frequencies in the interval $(f, f + df)$.

More formally, we consider the Fourier integral

$$S(f, T) = \int_0^T [N(t) - b] e^{-2\pi i f t} dt,$$

and we recall, for completeness, the relationship between $S(f, T)$ and the covariance function, $R(\tau)$, of $[N(t) - b]$. If

$$w(f) = \lim_{T \rightarrow \infty} \frac{2 |S(f, T)|^2}{T},$$

then

$$w(f) = 4 \int_0^\infty R(\tau) \cos 2\pi f \tau d\tau, \quad (11)$$

$$R(\tau) = \int_0^\infty w(f) \cos 2\pi f \tau df.$$

(Cf. Rice,⁴ p. 312 ff.)

At the same time, we have

$$\begin{aligned} R(\tau) &= E\{[N(t) - b][N(t + \tau) - b]\} \\ &= \psi(\tau) - b^2 \\ &= bg(\tau). \end{aligned}$$

Let $X(t)$ be any stochastic process which is known to be the occupancy of a telephone exchange of unlimited capacity, having a probability density of holding-time, and subject to Poisson traffic. From the preceding result it can be seen that the covariance function of $X(t)$ determines the distributions of the $X(t)$ process completely, since

$$\begin{aligned} a &= - \left. \frac{dR}{d\tau} \right]_{\tau=0}, \\ f(\tau) &= \int_\tau^\infty h(u) du = -a^{-1} \frac{dR}{d\tau}. \end{aligned}$$

If the holding-times are bounded by a constant, k , then readings of $N(t)$ taken further apart than k are uncorrelated. In fact, such values

are independent, because no call which contributes to $N(t)$ can survive until $(t + k)$, with probability 1.

Using (11), we see that

$$\begin{aligned} w(f) &= 4 \int_0^\infty \cos 2\pi f \tau R(\tau) d\tau \\ &= 4b \int_0^\infty \cos 2\pi f \tau g(\tau) d\tau \\ &= 4a \int_0^\infty \cos 2\pi f \tau \int_\tau^\infty \int_y^\infty h(u) du dy d\tau \quad (12) \\ &= \frac{2a}{\pi f} \int_0^\infty \sin 2\pi f \tau \int_\tau^\infty h(u) du d\tau \\ &= \frac{a}{\pi^2 f^2} \left[1 - \int_0^\infty \cos 2\pi f \tau h(\tau) d\tau \right]. \end{aligned}$$

Equation (12) expresses the mean square of the frequency spectrum of the fluctuations of the traffic away from the average in terms of the calling-rate and the cosine transform of the holding-time density, $h(u)$. The calling-rate appears only as a factor, and so does not affect the shape of $w(f)$. The function $w(f)$ is what Doob¹ (p. 522) calls the "spectral density function (real form)."

V EXAMPLE 1. $N(t)$ MARKOVIAN

Let the frequency $h(u)$ be negative exponential, so that

$$h(u) = \frac{1}{h} e^{-u/h}, \quad (13)$$

where h is the mean holding-time. It is shown in Riordan⁵ p. 1134, that $N(t)$ is Markovian if and only if $h(u)$ has the form (13). From page 523 of Doob¹ we know that the covariance function of a real, stationary Markov process (wide sense) has the form

$$R(\tau) = R(0)e^{-\alpha\tau}, \quad \alpha \text{ constant.} \quad (14)$$

Under the assumption (13), the covariance of $N(t)$ is

$$\begin{aligned} R(\tau) &= bg(\tau) = \frac{b}{h} \int_\tau^\infty \int_y^\infty h(u) du dy \\ &= b \int_\tau^\infty \int_y^\infty \frac{1}{h^2} e^{-u/h} du dy \\ &= be^{-\tau/h}, \end{aligned}$$

in agreement with (14). The spectral density can now be obtained from (11) or (12); it is

$$w(f) = \frac{4bh}{1 + 4\pi^2 f^2 h^2}.$$

This is the same as would be obtained for a Markov process that alternately assumed the values $+\sqrt{ah}$, $-\sqrt{ah}$ at the Poisson rate of $(2h)^{-1}$ changes of sign per sec. (Cf. Rice⁴ p. 325.)

VI EXAMPLE 2. HOLDING-TIME DISTRIBUTED UNIFORMLY IN (α, β)

Let $h(u)$ be constantly equal to $(\beta - \alpha)^{-1}$ in the interval (α, β) , and constantly 0 elsewhere. Then by (12),

$$\begin{aligned} w(f) &= \frac{a}{\pi^2 f^2} \left[1 - \frac{1}{\beta - \alpha} \int_{\alpha}^{\beta} \cos 2\pi f t \, dt \right] \\ &= \frac{a}{\pi^2 f^2} \left[1 - \frac{\sin 2\pi f \beta - \sin 2\pi f \alpha}{2\pi f (\beta - \alpha)} \right]. \end{aligned}$$

Now we see that

$$f(y) = \int_y^{\infty} h(u) \, du = \begin{cases} 1 & \text{for } y \leq \alpha \\ \frac{\beta - y}{\beta - \alpha} & \text{for } \alpha \leq y \leq \beta \\ 0 & \text{for } y \geq \beta \end{cases}$$

so that

$$R(\tau) = \begin{cases} a \left[\alpha - \tau + \frac{\beta - \alpha}{2} \right] & 0 \leq \tau \leq \alpha \\ \frac{a}{2} \frac{(\beta - \tau)^2}{\beta - \alpha} & \alpha \leq \tau \leq \beta \\ 0 & \tau \geq \beta \end{cases} \quad (15)$$

is the covariance function of the process $N(t)$ when holding-time is distributed uniformly in (α, β) .

If, formally, we let $(\beta - \alpha)$ approach 0 while keeping $\frac{1}{2}(\alpha + \beta)$ fixed, then the holding-times become concentrated in the neighborhood of the mean, h ; in the limit, as $h(u)$ tends to a singular normal distribution with mean, h , and variance zero, we obtain

$$w(f) = \frac{a}{\pi^2 f^2} [1 - \cos 2\pi f h] \quad (16)$$

as the spectral density function for the $N(t)$ process with constant holding-time, $h = \frac{1}{2}(\alpha + \beta)$. Similarly, from (15), we note that as the holding-times become singularly normal with mean, h , and variance zero, the covariance function becomes

$$R(\tau) = \begin{cases} = a(h - \tau) & 0 \leq \tau \leq h \\ = 0 & \tau \geq h. \end{cases}$$

We can express (16) as

$$w(f) = 2ah^2 \left(\frac{\sin \pi fh}{\pi fh} \right)^2,$$

and note that this is exactly like the power spectrum of a random telegraph wave constructed by choosing values $+\sqrt{ah}$, $-\sqrt{ah}$ with equal probability and independently for each interval of length, h . (Cf. Rice,⁴ page 327.)

REFERENCES

1. J. L. Doob, *Stochastic Processes*, John Wiley and Sons, New York, 1953.
2. C. Palm, *Intensitätsschwankungen im Fernspreckverkehr*, Ericsson Technics, **44**, pp. 1-189, 1943.
3. F. W. Rabe, *Variations of Telephone Traffic*, Elec. Commun., **26**, pp. 243-248, 1949.
4. S. O. Rice, *Mathematical Analysis of Random Noise*, B.S.T.J., **23**, pp. 282-332, 1944, and **24**, pp. 46-156, 1945.
5. J. Riordan, *Telephone Traffic Time Averages*, B.S.T.J., **30**, pp. 1129-1144, 1951.

High-Voltage Conductivity-Modulated Silicon Rectifier

By H. S. VELORIC and M. B. PRINCE

(Manuscript received May 1, 1957)

Silicon power rectifiers have been made which have reverse breakdown voltages as high as 2,000 volts and forward characteristics comparable to those obtained in much lower voltage devices. It is shown that the magnitude and temperature dependence of the currents can be explained on the basis of space-charge generated current with a trapping level 0.5 eV below the conduction band or above the valence band. The breakdown voltage of a P^+IN^+ diode is computed from avalanche multiplication theory and is shown to be a function of the width of the nearly intrinsic region. A simple diffusion process is evaluated and shown to be adequate for diode fabrication. The characteristics of devices fabricated from high-resistivity compensated, floating-zone refined, and gold-doped silicon are presented. The surface limitation to high inverse voltage rectifiers is discussed.

I INTRODUCTION

The desire for high voltage rectifiers in the electronic industry has pushed the peak inverse voltage of solid state rectifiers to higher and higher values. The purpose of this paper is to present some of the considerations necessary in designing a device with a high inverse voltage and an excellent forward characteristic. In many cases the device characteristics are predictable. Conversely, high voltage diodes are excellent tools for studying many solid state phenomena.

It has been shown¹ that it is possible by the use of the conductivity modulation principle to separate the design of the forward current-voltage characteristic from the reverse current-voltage characteristic of a silicon p - n junction rectifier. Units have been fabricated by the diffusion of boron and phosphorus into high resistivity material, that have reverse breakdown voltages in the range of 1,000 to 2,000 volts.

The reverse currents are of the order of a microampere per square centimeter at room temperature and increase approximately as the square

root of the applied voltage. The magnitude, voltage dependence, and temperature dependence of the reverse currents can be explained as due to space-charge generated current² with a trapping level 0.5 eV from either the conduction or valence band. These effects will be discussed in Section II.

In Section III the breakdown voltage and its dependence on the resistivity and width of the high resistivity region of the rectifier will be considered.

In the next section the forward current is discussed and explained by considering both a space-charge region generated current and a diffusion current that takes into account high levels of minority carrier injection.³

Device processing information is given in Section V, together with an evaluation of different sources of high resistivity silicon. The devices to be discussed in this paper have been processed with high resistivity *p*-type material, although some devices have been made with *n*-type material.

Finally, a discussion of some surface problems associated with high voltage rectifiers is given in Section VI.

Although this paper is entitled "High-Voltage Conductivity-Modulated Silicon Rectifier", the theoretical arguments are applicable to all semiconductor diodes. However, the experimental results have been limited by considering only high voltage diodes.

II REVERSE CURRENT-VOLTAGE CHARACTERISTIC

2.1 Theory

The simple theory⁴ for a *p-n* junction yields an expression for the reverse saturation current density (I_0) which is:

$$I_0 = q \left[n_p \left(\frac{D_n}{\tau_n} \right)^{1/2} + p_n \left(\frac{D_p}{\tau_p} \right)^{1/2} \right], \quad (2-1)$$

where q is the electron charge, n_p is the equilibrium electron density in *p*-type material, p_n is the equilibrium hole density in *n*-type material, D_n and D_p are the diffusion constants for electrons and holes, and τ_n and τ_p are the minority carrier lifetimes for electrons and holes.

When reasonable numbers are substituted into (2-1), I_0 at room temperature is of the order of 10^{-10} amperes per square centimeter. This quantity doubles with every increase of 4° C. The theory also contains no voltage dependence of this current. Even when breakdown multiplication⁵ is taken into account, there is essentially no voltage dependence

at voltages less than half the breakdown voltage. The magnitude and temperature and voltage dependences of measured diodes do not agree with these theoretical values at room temperatures.

Recently, Pell⁶ has shown that the reverse currents at low temperatures in germanium, and at room temperatures in silicon, are dominated by space-charge generated current. The space-charge generated current density (I_{sc}) is given by

$$I_{sc} = q W G M, \quad (2-2)$$

where W is the width of space-charge region, G is the generation rate of hole-electron pairs in the space-charge region, and M is the breakdown multiplication ($M \sim 1$ except near the breakdown voltage). G is given by²

$$G = \frac{n_1 p_1}{\tau_{p0} n_1 + \tau_{n0} p_1}, \quad (2-3)$$

where n_1 and p_1 are the densities of electrons and holes respectively if the Fermi levels were at the energy level of the recombination centers, and τ_{n0} and τ_{p0} are the minority carrier lifetimes of electrons and holes respectively in heavily doped p -type and n -type silicon. This expression assumes constant generation over the space-charge region. Thus,

$$n_1 = N_c \exp \frac{q}{kT} (V_r - V_c) = n_i \exp \beta (V_r - V_i), \quad (2-4a)$$

and

$$p_1 = N_v \exp \frac{q}{kT} (V_v - V_r) = n_i \exp -\beta (V_r - V_i), \quad (2-4b)$$

where V_r is the recombination level above the valence band edge V_v , V_i is the midband intrinsic level, $\beta = q/kT$, N_c and N_v are the effective densities of states in the conduction and valence bands $\cong 2.4 \times 10^{19} (T/300)^3$, V_c is the conduction band edge, k is the Boltzmann's constant, and T is the absolute temperature.

Substituting (2-4) into (2-3), one obtains:

$$G = \frac{n_1}{2\sqrt{\tau_{n0}\tau_{p0}}} \frac{1}{\cosh \left[\beta (V_r - V_i) + \frac{1}{2} \ln \frac{\tau_{p0}}{\tau_{n0}} \right]}. \quad (2-5)$$

For the diffused silicon junctions under consideration, it has been found⁷ that τ_{n0} equals 1.2×10^{-6} seconds and τ_{p0} equals 0.4×10^{-6} seconds. Also, $n_i = 3.74 \times 10^{15} T^{3/2} e^{-6250/T}$ and $V_i = 0.54$ volts. Using these

numbers, (2-6) becomes:

$$G = \frac{1.25 \times 10^{16} \left(\frac{T}{300}\right)^{3/2} e^{20.8(1-300/T)}}{\cosh \left[38.62 \left(\frac{300}{T}\right) (V_r - 0.54) - 0.55 \right]}, \quad (2-6a)$$

and

$$G = G_{300}(V_r)f(T, V_r), \quad (2-6b)$$

where

$$G_{300}(V_r) = \frac{1.25 \times 10^{16}}{\cosh [38.62(V_r - .54) - 0.55]}, \quad (2-7a)$$

and

$$f(T, V_r) = \left(\frac{T}{300}\right)^{3/2} e^{20.8(1-300/T)} \frac{\cosh [38.62(V_r - 0.54) - 0.55]}{\cosh \left[38.62 \left(\frac{300}{T}\right) (V_r - 0.54) - 0.55 \right]}. \quad (2-7b)$$

In (2-6b), $G_{300}(V_r)$ is the generation rate for a recombination level at V_r equal to 300° K, and $f(T, V_r)$ is the temperature variation of G for a recombination level at V_r normalized to 300° K.

Curves of $f(T, V_r)$ are given in Fig. 1 for several values of V_r with a curve $g(T)$ which is the temperature variation of the reverse saturation current (I_0). Table I gives values for G_{300} for various V_r .

In the reverse biased diffused junctions made with high resistivity material, the junction may be considered abrupt. Therefore, the width (W) of the space-charge region, when the junction is reverse biased to a voltage V , is given by

$$W = \left(\frac{\kappa V}{2qN_A} \right)^{1/2} = 3.14 \times 10^{-5} (V\rho_p)^{1/2} \text{ cm} \quad (2-8)$$

where the units after the first equal sign are electrostatic, and κ is the dielectric constant. In the second expression, V is in volts, and ρ_p , the base material resistivity, expressed in ohm-centimeters. Thus,

$$I_{sc} = 4 \times 10^{-24} G_{300}(V_r)f(T, V_r)[V\rho_p]^{1/2} \text{ amperes-cm}^{-2}. \quad (2-9)$$

It is seen that I_{sc} varies theoretically as the square root of the reverse voltage for values of V less than $\frac{1}{2} V_B$, the breakdown voltage, in which range avalanche multiplication is negligible. The quantity I_{sc} varies inversely with $N_A^{1/2}$ and will be large for high voltage devices with small N_A . The I_{sc} at 300° K for a rectifier with 40 ohm-centimeter base ma-

terial and a reverse bias of 100 volts is given in Table I as a function of V_r . The numbers compare with 8×10^{-10} ampere per square centimeter for I_0 . Thus, from diode measurements at room temperature and above, one could not observe V_T less than 0.3 eV from either the conduction or valence band. In fact, from a measurement of the temperature dependence of the reverse currents, one can determine only the recombination level lying closest to the center of the forbidden band. This can be seen more clearly from the following argument: There will be a contribution to the reverse current from the diffusion current $I_{0300}g(T)$ which varies with temperature as $g(T)$. There will be contributions to the reverse cur-

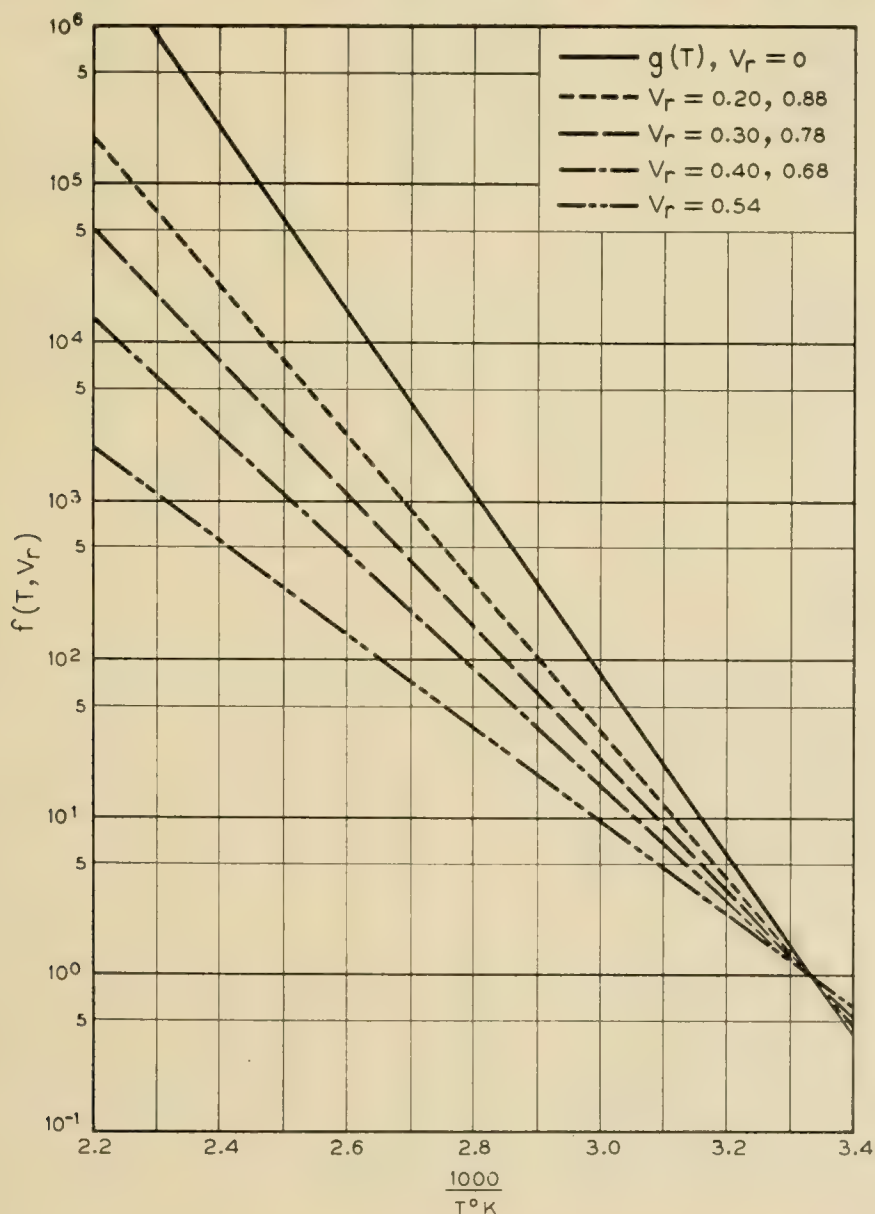


Fig. 1 — The temperature variation of the generation rate, $f(T, V_r)$, for several values of the recombination level, V_r .

TABLE I — VALUES OF G AND SPACE CHARGE GENERATED CURRENT AT 300° K FOR VARIOUS VALUES OF THE TRAPPING LEVEL V_r

$$T = 300^\circ \text{ K} \quad E_g = 1.08 \text{ eV} \quad np = 3 \times 10^{20} \text{ cm}^{-6}$$

$$\tau_{n0} = 1.2 \times 10^{-6} \text{ sec} \quad \tau_{p0} = 0.4 \times 10^{-6} \text{ sec}$$

V_r Volts above Valence Band	$G_{300} \text{ cm}^{-3} \text{ sec}^{-1}$	$I_{sc} (V = -100 \text{ volts}, \rho = 40 \text{ ohm-cm})$ microamperes/cm ²
0.10	5.92×10^8	1.88×10^{-7}
0.20	2.84×10^{10}	9.03×10^{-6}
0.30	1.35×10^{12}	4.29×10^{-4}
0.40	6.5×10^{13}	2.06×10^{-2}
0.50	3.02×10^{15}	0.96
0.54	1.08×10^{16}	3.43
0.58	8.1×10^{15}	2.58
0.68	1.95×10^{14}	6.2×10^{-2}
0.78	3.96×10^{12}	1.26×10^{-3}
0.88	8.45×10^{10}	2.69×10^{-5}
0.98	1.78×10^9	3.75×10^{-7}

rent by the individual trapping centers given by $I_{sc300}(V_r)f(T, V_r)$, where $I_{sc300}(V_r)$ is the current at 300° K due to generation at recombination centers located at the level V_r , and $f(T, V_r)$ is the temperature variation of the generation rate. Thus, the total reverse current is given by

$$I_{\text{reverse}} = I_{o300}g(T) + \sum_{V_r} I_{sc300}(V_r)f(T, V_r), \quad (2-10)$$

where the summation is over all recombination levels. The relative currents at 300° K are given in Table I. The greatest contribution at 300° K is due to the level nearest the center of the forbidden band. As the temperature increases, all the terms under the summation sign approach each other. Before a second recombination level contributes significantly to the reverse current, however, the saturation current will have become the most important component.

2.2 Experimental Results

To evaluate the theory for the reverse currents in silicon N⁺P junctions, careful measurements were made on five typical units for the reverse current-voltage characteristics at various temperatures from 300° K to 435° K. The curves were taken with a X-Y recorder. The voltage ranged from 0 to 200 volts so that multiplication effects were completely negligible.

From the recorded data, curves of I_r versus $1,000/T^\circ \text{ K}$ were plotted for $V = -10, -40$, and -160 volts. The set of curves for diode No. 3

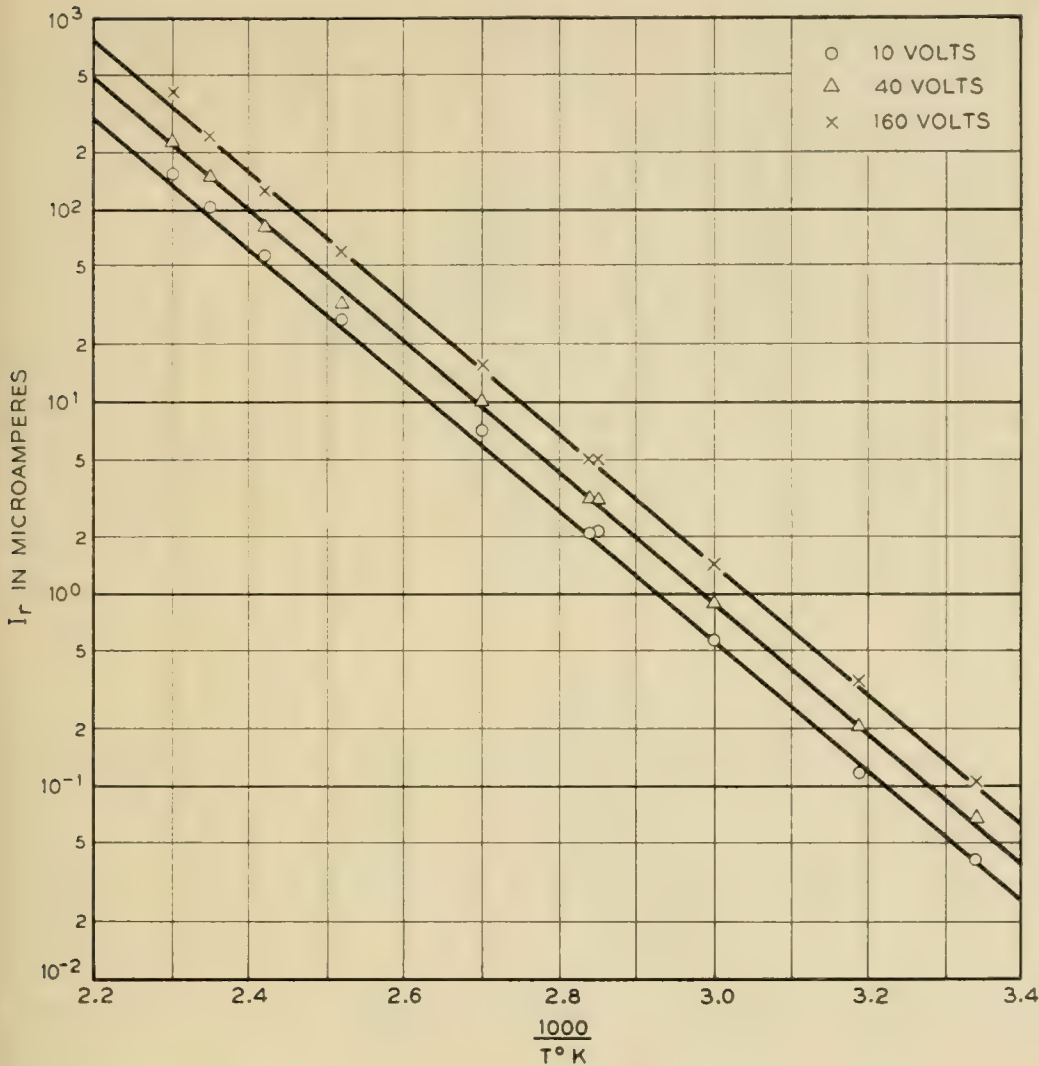


Fig. 2 — The temperature variation of “reverse current” for a typical diode at -10 , -40 , and -160 volts.

is given in Fig. 2. The slope of these curves indicates that the recombination level lies near 0.5 eV below the conduction band or above the valence band. The junction area of this device is 0.015 cm²; thus, the current density at 300° K and at -100 volts is 4.4 microamperes per square centimeter. This compares with the order of one microampere as listed in Table I. This suggests that the $(\tau_{n0}\tau_{p0})^{\frac{1}{2}}$ is overestimated. The agreement of this measurement with theory is reasonable.

The voltage variation of the reverse currents does not agree with theory as well as the magnitude and temperature dependence. The experimental results give, as the voltage dependence, an expression:

$$I_r \sim V^{1/N},$$

where N equals 2.9 . This compares with the theoretical value of $N = 2$.

Some of this discrepancy can be attributed to the fact that the junction is not truly an abrupt junction. A "graded" junction would yield N equals 3. Measurements of capacitance versus voltage, which essentially measure the width of the space-charge region, yield N equals 2.4. Thus, these devices in the relatively low voltage range still have some contribution from the gradation of the diffused junction.

The highest temperature points in Fig. 2 deviate above the straight lines. This deviation can be attributed to the onset of the contribution from the I_0 component. Calculations indicate that, by $T = 220^\circ \text{C}$, the contribution to the reverse current by the space-charge current is equaled by the saturation current and that, by $T = 320^\circ \text{C}$, the space-charge-generated current is negligible compared to the saturation current.

III BREAKDOWN VOLTAGE OF PN AND PIN JUNCTIONS

3.1 Theory

It has been demonstrated that, in germanium^{8, 9} and silicon, reverse biased junctions breakdown as a result of a solid state analogue of the Townsend β Avalanche Theory. Multiplication and breakdown occur when electrons or holes are accelerated to energies sufficient to create hole-electron pairs by collisions with valence electrons. The breakdown phenomena in silicon for graded and step junctions has been previously considered.^{8, 10} Depending on the impurity distribution, the field in the junction will be a function of distance and will have a maximum value in the region of zero net impurity concentration. The breakdown voltage is a critical function of the space-charge distribution.

In this section the existing multiplication theory is extended to the case of PIN junctions. It is shown that relatively wide intrinsic regions are required to obtain breakdown voltages greater than 1000 volts.

Fig. 3 is a plot of the impurity, charge, and field distributions in PIN and $P\pi N$ junctions. Fig. 3(a) schematically illustrates the geometry of the three region devices considered, and Fig. 3(b) is a plot of the impurity distribution. In this analysis step junctions will be assumed. For the PIN junction there are no uncompensated impurities in the intrinsic region, and no net charge. At low reverse voltage, the field will sweep through the intrinsic layer and will increase with increasing reverse bias until the breakdown field is reached.

Absolutely intrinsic material is not yet available, and devices are made from high resistivity π -type material. In this class of devices there is some uncompensated impurity and charge in the center region. The field will have a maximum value at the $N^+ \pi$ junction and will decrease

with increasing distance into the π region. At sufficiently high reverse bias the field may sweep into the P^+ region.

Breakdown in silicon⁸ is a multiplicative process described by

$$1 - \frac{1}{M} \int_0^W \alpha_1 dx, \quad (3-1)$$

where M is the multiplication factor, W is the space-charge width, and α_1 is the rate of ionization which is a strong function of the field in the junction. For a PIN structure, the field is constant, at breakdown M approaches ∞ , and

$$\alpha_1 W = 1. \quad (3-2)$$

The ionization rate at breakdown is then a simple function of the width of the intrinsic region. McKay⁸ and Wolf¹¹ have considered α_1 as a func-

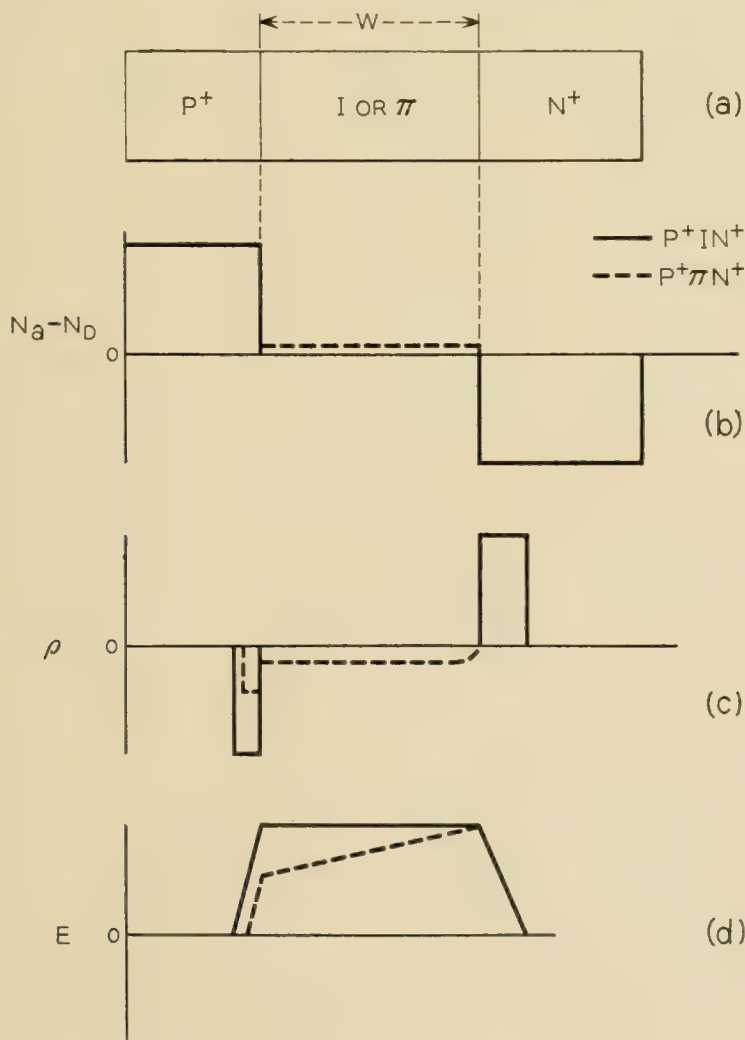


Fig. 3 — Impurity, charge and field distribution in PIN and P π N junctions.

tion of the field. If α_1 is fixed, the field at breakdown can be determined. The breakdown voltage is $\int_0^w E dx$.

Fig. 4 is a plot of breakdown voltage as a function of space-charge width for PN and PIN diodes. The PIN values are calculated; the PN data is previously unpublished data supplied by K. G. McKay.

Some interesting observations can be made from Fig. 4:

1. The plot of breakdown voltage versus barrier width for a PN step junction assumes that the space-charge region does not extend through the high resistivity side of the junction. For this class of junctions the breakdown voltage is determined by the impurity concentration as shown in Fig. 5. The plot of breakdown versus space-charge width for a PIN diode assumes that the space-charge region extends from the P to N region at very low bias, and that it is limited by the width of the I region. If a constant field is assumed in the I region, the breakdown voltage is a function of the barrier width.

2. Although the space-charge region can reach through the I region at low bias, the avalanche breakdown voltage is a function of the width of the I region.

3. For the devices considered here with π regions in the order of 10^{-2} cm, the maximum breakdown voltage is in the order of 2,000 volts.

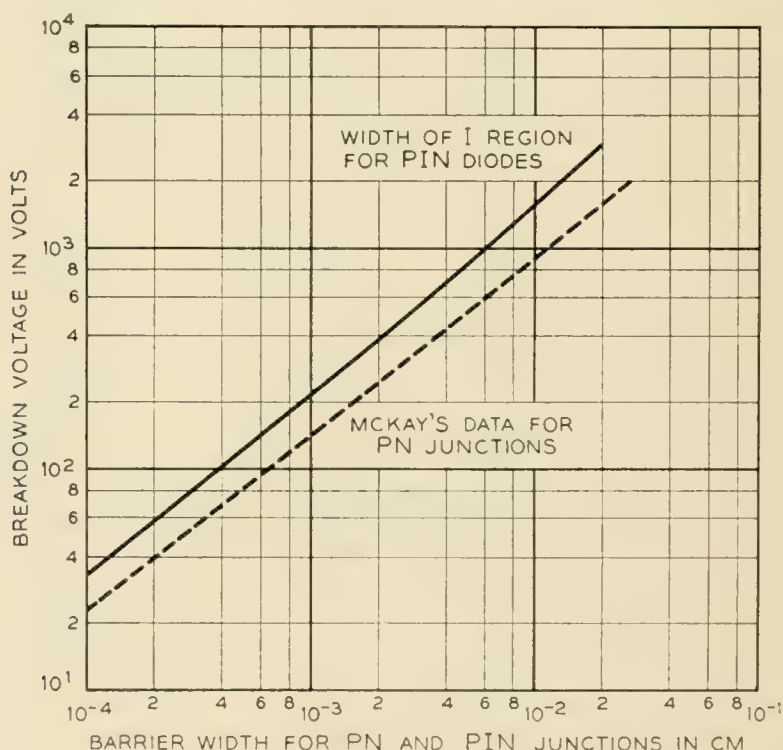


Fig. 4 — Breakdown voltage as a function of barrier width for PN and PIN junctions.

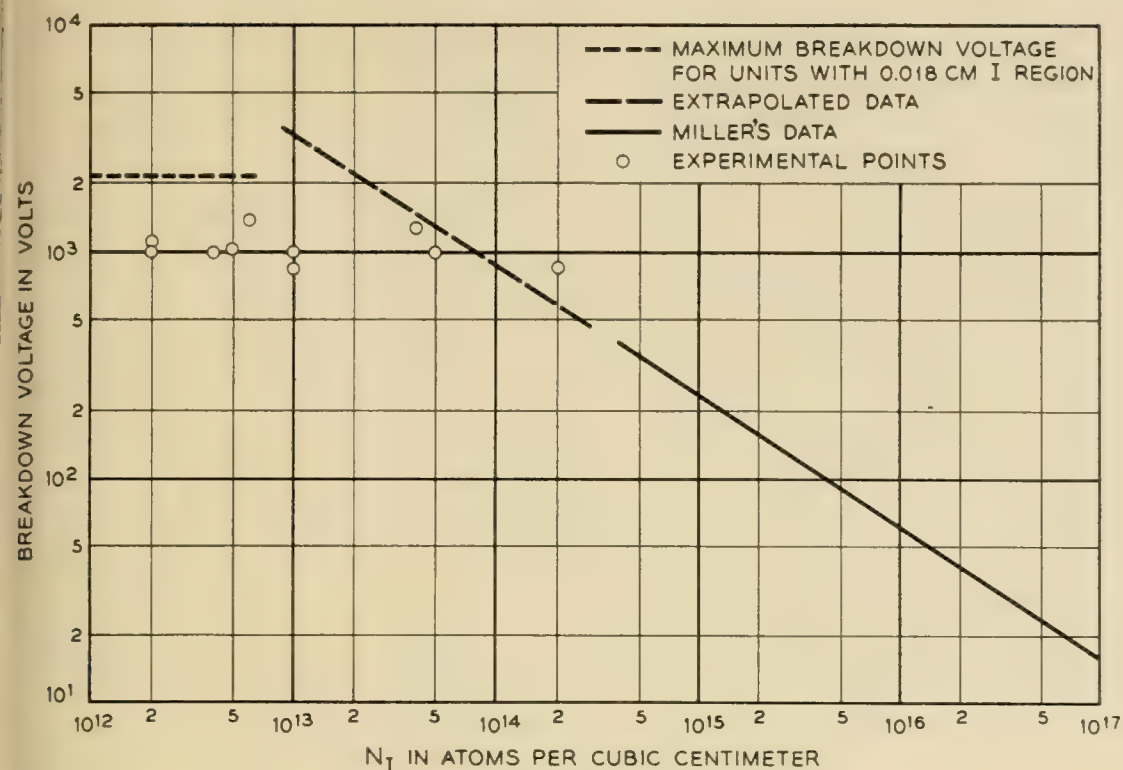


Fig. 5 — Breakdown voltage versus impurity concentration for silicon step junctions.

3.2 Experiment

Fig. 5 is a plot of breakdown voltage versus impurity concentration for silicon step junctions. The plot above 300 volts is extrapolated from the data of Miller¹⁰ and Wilson.⁹

Capacity data, discussed in Section V, indicates that many devices show body breakdown. A few rectifiers break down at voltages as high as 2,000 volts. In many high voltage devices the breakdown voltage is not limited by geometry but by surface problems.

IV FORWARD CURRENT-VOLTAGE CHARACTERISTIC

4.1 Theory

It will be shown in this section that the forward current-voltage characteristic as well as the reverse characteristic can be completely explained by considering both a space-charge region generated current and a diffusion current. The diffusion current component must also take into consideration the effect of high injection levels of minority carriers.

According to the Shockley-Read² theory, the rate of recombination, U , of holes and electrons in a semiconductor is given by:

$$U = -G = \frac{pn - n_i^2}{\tau_{p0}(n + n_1) + \tau_{n0}(p + p_1)} \quad (4-1)$$

where p , and n are the instantaneous concentrations of holes and electrons, respectively. When a PN junction is forward biased, holes and electrons are injected into the space-charge region which has been reduced in width. Some of these carriers diffuse through the space charge region and give rise to the normal diffusion current when the excess minority carriers recombine with majority carriers in field free regions. The other carriers recombine according to (4-1) in the space-charge region giving rise to what is called the space-charge generated current. In the reverse biased junction, the current is due to carriers generated in the space-charge region; whereas, in the forward biased junction, the current is due to recombination of carriers. The quantity U is large in the space-charge region since both p and n are large in this region. In the field free regions, however, one of these quantities is usually small and the product deviates only slightly from n_i^2 .

The space-charge generated current, I_{sc} , is given approximately by:¹²

$$I_{sc} = \frac{2qWn_i}{(\tau_{n0}\tau_{p0})^{1/2}} \frac{\sinh \beta \frac{V}{2}}{\beta(V_B - V)} f(b), \quad (4-2)$$

where V_B is the built-in potential of the junction, and $f(b)$ is discussed in Reference 12 and is approximately 1.5 for recombination centers near the intrinsic level as is the case for the diodes under consideration. For shallower recombination levels the function $f(b)$ is much smaller and depends strongly upon the forward applied voltage.

For the forward-biased junction, the space-charge region is narrow, the concentration gradient can be considered linear and W is given by the following expression:⁴

$$W = 4.35 \times 10^2 \left(\frac{V_{\text{junction}}}{\alpha} \right)^{1/3} \text{ cm}, \quad (4-3)$$

where V_{junction} is the total potential across the junction in volts and α is the concentration gradient at the junction in cm^{-4} . These are given by:

$$\begin{aligned} V_{\text{junction}} &= V_{\text{built-in}} - V \\ &= kT/q \ln (N_A N_D / n_i^2) - V \\ &= 0.792 - V. \end{aligned} \quad (4-4)$$

$$\text{Also, } \alpha = \frac{C_0}{\sqrt{\pi D t}} e^{-x_j^2 / 4 D t} \text{ for diffused junctions,} \quad (4-5)$$

where C_0 = surface concentration of diffusant = $3 \times 10^{19} \text{ cm}^{-3}$,

D = diffusion constant = 3×10^{-12} cm²/sec,

t = diffusion time = 5.7×10^4 sec,

x_j = junction depth below surface = 0.003 cm.

When these numbers are substituted into the equations, at 300° K:

$$W = 9.25 \times 10^{-4} (0.792 - V)^{1/3} \text{ cm.} \quad (4-6)$$

For the diodes under consideration:

$$\tau_{n0} = 1.2 \times 10^{-6} \text{ sec,} \quad \tau_{p0} = 0.4 \times 10^{-6} \text{ sec.}$$

When these expressions are substituted into (4-2), one obtains at 300° C:

$$I_{sc} = 2.8 \times 10^{-7} \frac{\sinh 19.31 V}{(0.792 - V)^{2/3}} \text{ amp/cm}^2. \quad (4-7)$$

In order to fit the experimental data, it is necessary to multiply (4-7) by a factor of 5. This may be due to an overestimation of $(\tau_{n0} \tau_{p0})^{1/2}$. Therefore, the equation which shall be used in the remainder of this section will be:

$$I_{sc} = 1.4 \times 10^{-6} \frac{\sinh 19.31 V}{(0.792 - V)^{2/3}} \text{ amp/cm}^2. \quad (4-8)$$

A plot of this expression is given in Fig. 6.

The normal diffusion current for low level diffusion,⁴ I_{DL} , is given by

$$I_{DL} = I_0 (e^{qV/kT} - 1) \quad (4-9)$$

where I_0 is given by (2-1). I_0 for the diodes under discussion is approximately 8×10^{-10} ampere/cm² at 300° K. When the injected minority carrier density approaches the equilibrium majority carrier density, the form of (4-9) changes. The high injection level diffusion current, I_{DH} , is then given by³

$$I_{DH} = I_{DH0} (e^{qV/2kT} - 1), \quad (4-10)$$

where I_{DH0} equals $qn_i s / \tau$, and s equals the width of the high resistivity region. For the diodes under discussion, I_{DH0} is approximately 2×10^{-6} amperes/cm² at 300° K. A current-voltage plot of these currents at 300° K for $V_r = 0.50$ is given in Fig. 6 together with their sum. It can be observed that the resulting characteristic starts with slope of qV/kT and bends over to a slope of $qV/2kT$ near 0.10 volt. The slope increases again to near qV/kT at 0.35 volts and decreases once more to $qV/2kT$ above 0.40 volts giving a bump to the over-all characteristic.

Next, consider the temperature dependence of the coefficients of the forward current components

$$I_{sc0} \sim n_i(T), \quad (4-11a)$$

$$I_0 \sim n_i^2 \frac{D_n(T)^{1/2}}{\tau_n(T)} \sim n_i(T)^2, \quad (4-11b)$$

and

$$I_{DH0} \sim \frac{n_i(T)}{\tau(T)} \sim n_i(T). \quad (4-11c)$$

The largest variation of these coefficients is due to the variation of

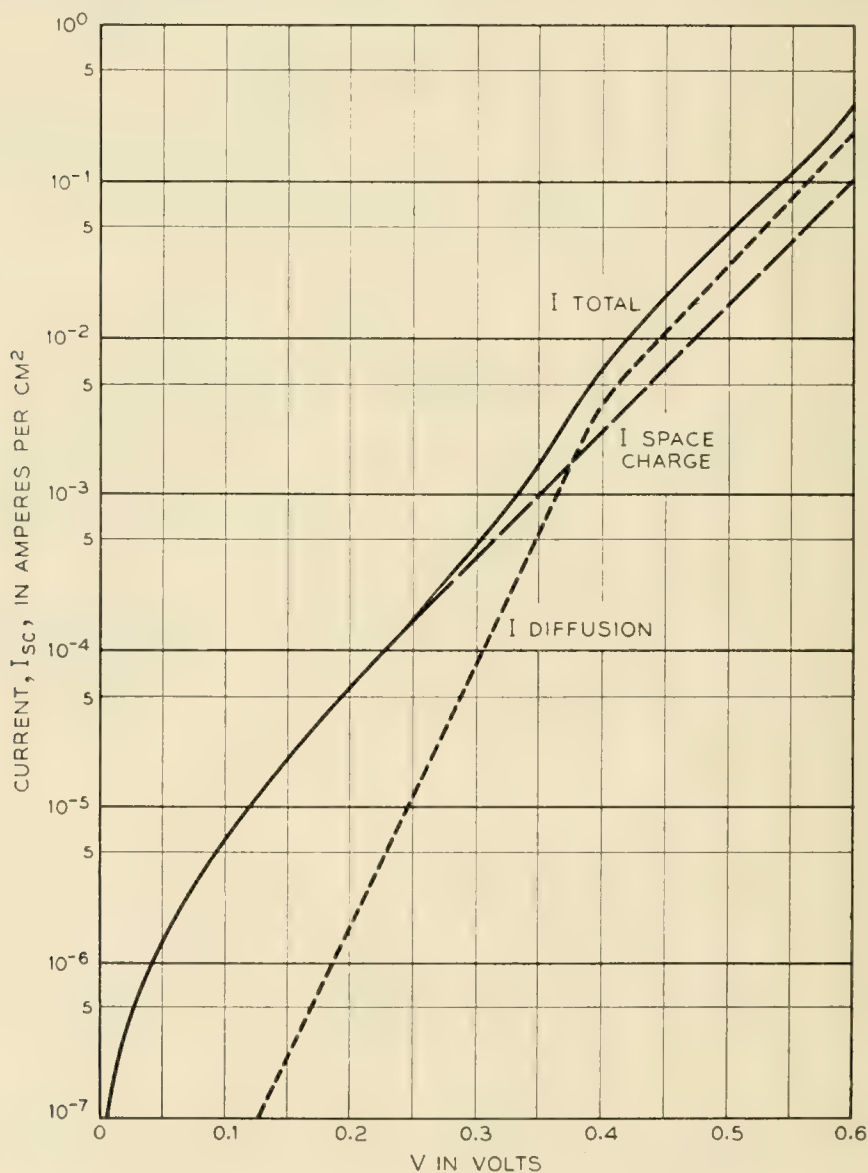


Fig. 6 — The two components of current for a forward biased junction.

n_i with temperature. Fig. 7 gives a plot of this variation. The temperature variations of the other parameters are all small compared to that of n_i . Thus, as in the case of the reverse currents, at sufficiently high temperatures, the diffusion current makes the more important contribution.

In the case of the forward current, I_{sc} is relatively insensitive to the distribution of impurities; therefore, the results of this section are important for all forward-biased diodes. In high-voltage diodes, to keep the resistive voltage drop small, it is necessary to maintain high minority carrier lifetime in the center region. The diffusion length of injected

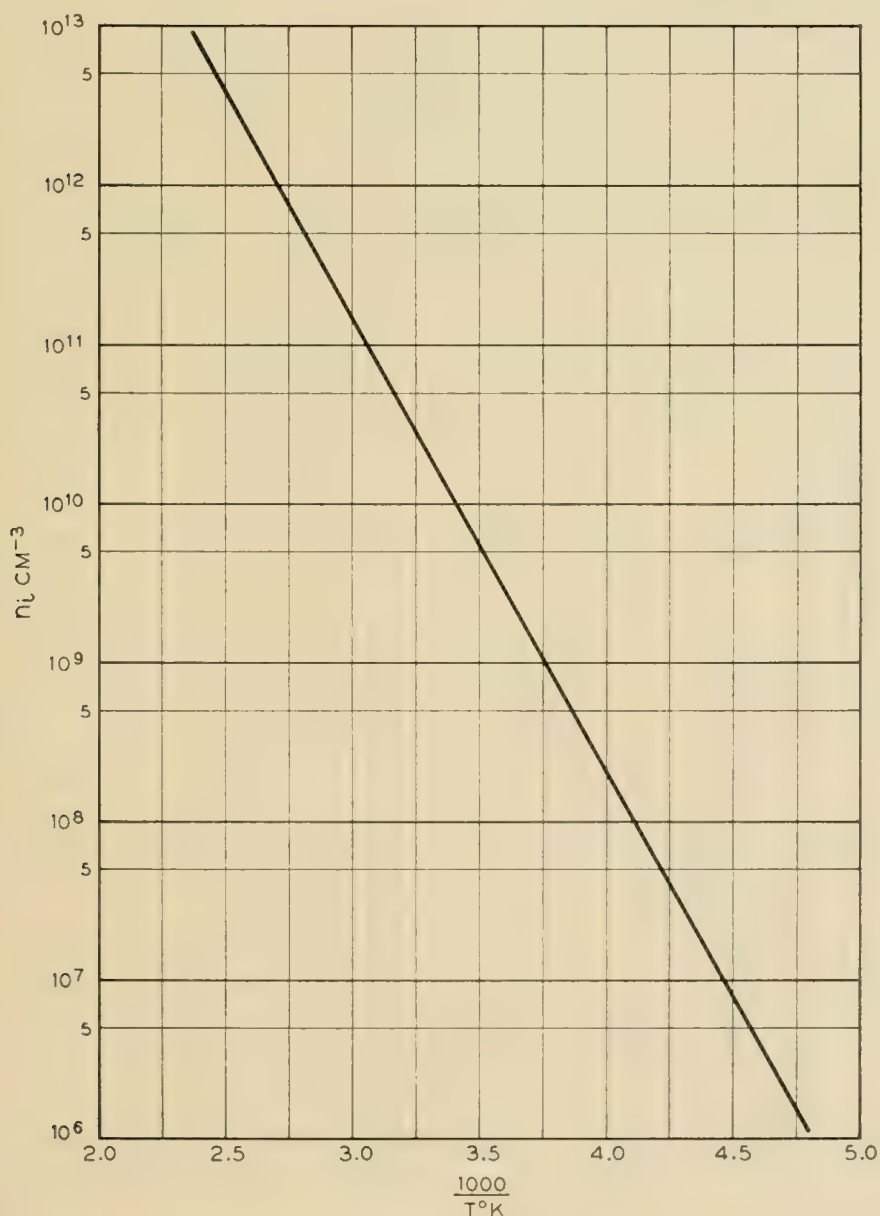


Fig. 7 — The variation of n_i with temperature.

minority carriers should be the order of or larger than the center region width.

4.2 Experimental Results

The forward characteristics of the five typical high voltage rectifiers mentioned in Section 2.2 were measured with a X-Y recorder, and all showed similar shapes. Diode No. 3 will be discussed in detail in this section.

The forward current-voltage characteristic was measured at three temperatures: 220° K, 300° K, and 375° K. Measurements below cur-

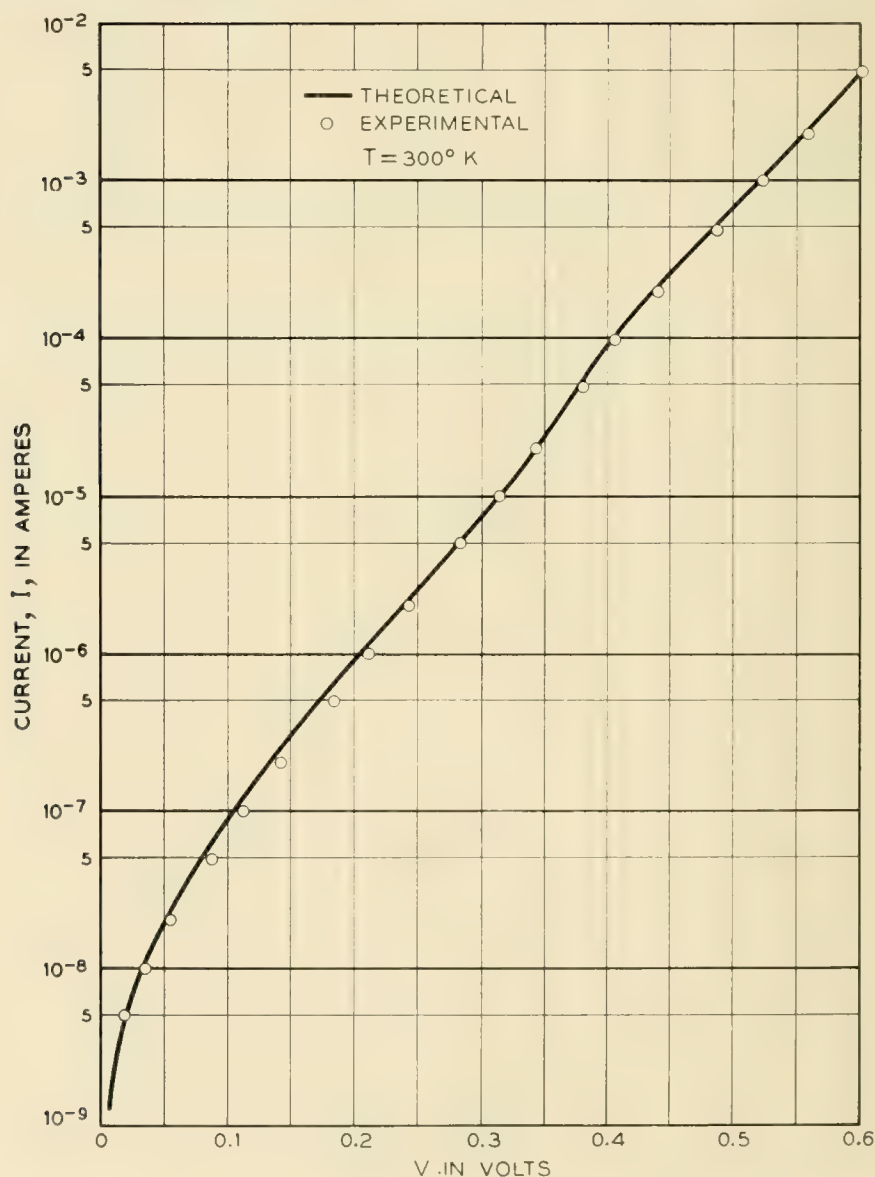


Fig. 8 — The calculated and observed current-voltage characteristic of a forward biased junction at 300° K.

rents of one microampere were made only at 300° K. Currents above 10 milliamperes were not measured since internal power losses would cause temperature variations. The unit has a junction area of 0.015 cm², junction lifetime of 4 microseconds at 300° K, S equal to 0.008 cm, and N_A equal to 3×10^{14} cm⁻³. When these numbers are substituted into the expressions for the coefficients, one obtains, at 300° K,

$$I_{sc0} = 1.4 \times 10^{-8} \text{ amperes,}$$

$$I_0 = 1.2 \times 10^{-11} \text{ amperes,}$$

$$I_{DH0} = 3.0 \times 10^{-8} \text{ amperes.}$$

Fig. 8 shows a semilogarithmic plot of the current-voltage characteristic at 300° K over a range of 6½ decades. The circles represent measured points and the solid line is the theoretical curve.

Using the variation of n_i given in Fig. 8 and the temperature variations as given in (4-11), one can obtain the coefficients for any temperature. This has been done for two temperatures, 220° K and 375° K. Figs. 9 and 10 show the theoretical and experimental plots at 375° K and 220° K

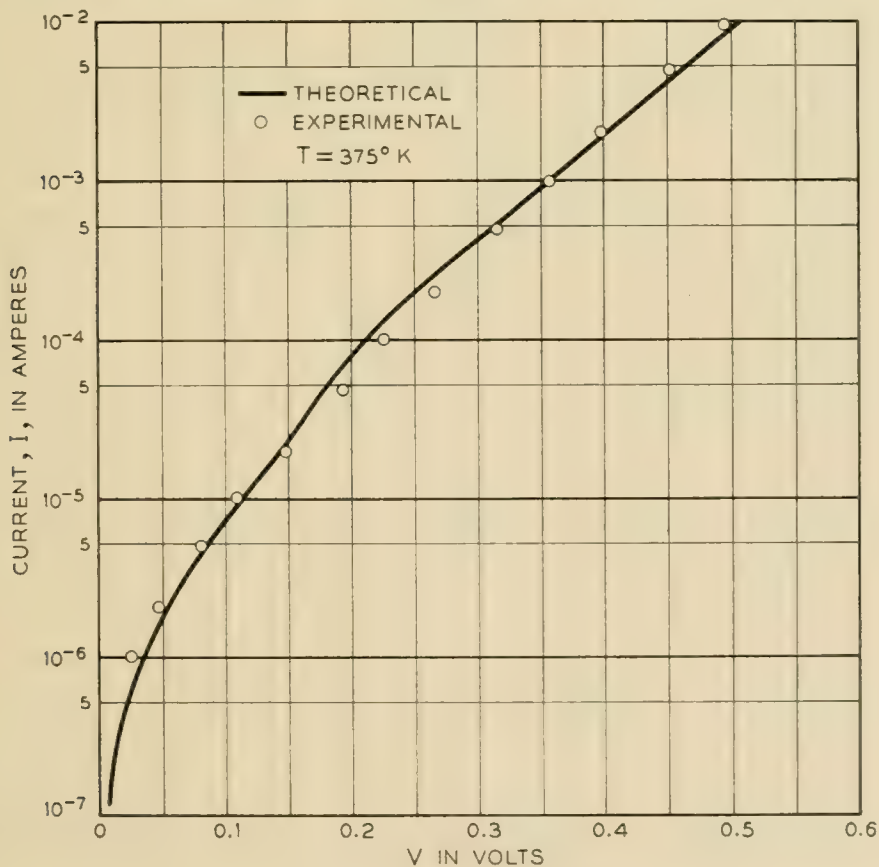


Fig. 9 — The calculated and observed current-voltage characteristic of a forward biased junction at 375° K.

respectively. The circles represent measured points and the solid lines are the calculated theoretical curves. It is observed that the fit in Fig. 9 is quite good; whereas, the fit in Fig. 10 is not as good as at the other temperatures. However, even this figure shows good qualitative agreement of the deviation from a straight line. Some of the factor of two discrepancy in Fig. 10 can be ascribed to the temperature variation of the other parameters, and some to a possible error in the measurement of temperature which would be reflected in the value of n_i .

It should be noted that at all temperatures the IR drop in the high resistivity region is not observable to the limits of the experimental measurements of forward current, 10 milliamperes. This is due to the fact that the region has been conductivity modulated by the forward current. This requires a sufficient minority carrier lifetime in the region so that most of the injected carriers diffuse across the region before recombining. Such lifetimes can be maintained in diffused junctions⁷

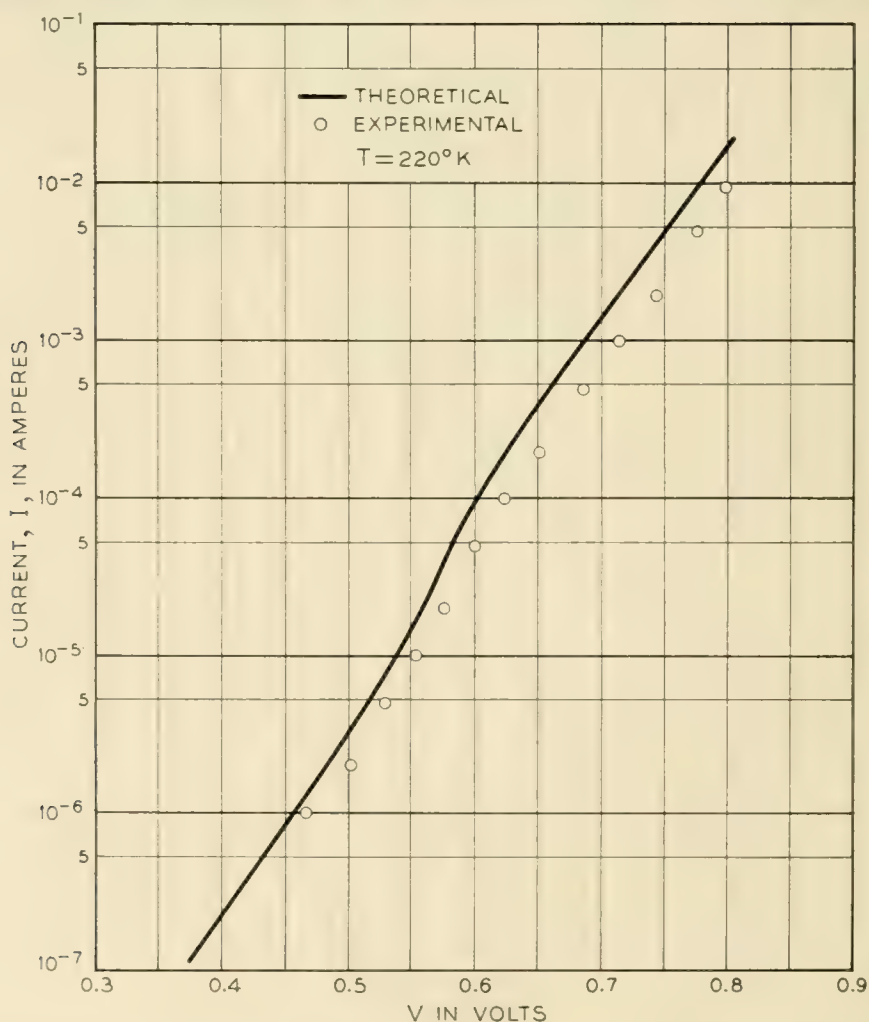


Fig. 10 — The calculated and observed current-voltage characteristic of a forward biased junction at 220° K.

to permit the high resistivity region to be at least as wide as 0.025 centimeters. Thus even in high voltage rectifiers it is still possible to design the forward and reverse current voltage characteristics independently.

V DEVICE PROCESSING

5.1 Silicon Material

Fig. 5 shows that step junctions which break down at over a thousand volts must have a background impurity concentration $\leq 10^{14}$ atoms/cm³. The highest grade commercial semiconductor silicon has 5×10^{14} impurities/cm³ (20–50 Ω cm P type). This material must be processed to reduce the impurity level. To date, high voltage devices have been processed from four types of high resistivity material: floating zone refined, compensated, gold diffused, and horizontal zone refined silicon.

Some silicon was prepared by adding N-type impurities to reduce $|N_D - N_A| < 10^{14}$. Maintaining this delicate balance in material where $N_D \simeq N_A$ is difficult. The boron is relatively uniformly distributed since the distribution constant is close to unity. N-type impurities are less uniformly distributed in the crystal since the distribution constants are considerably less than unity. High resistivity compensated silicon is full of N- and P-region striations. The units processed from this material generally had poor electrical characteristics.

Table II is a typical contour of a compensated crystal. The resistivity varies around the crystal and changes along the length of the crystal. At the bottom of the crystal the resistivity goes through a maximum. The tail end is converted from P to N type.

A number of devices have been fabricated from silicon processed with

TABLE II — A TYPICAL CONTOUR OF A HIGH
RESISTIVITY COMPENSATED CRYSTAL
Crystal A-161, Oriented 111, Rotated $\frac{1}{2}$ RPM

Distance from seed (inches)	Resistivity (Ω cm) at Angle				Impurity Type
	0°	90°	180°	270°	
$\frac{1}{2}$	28	33	23	30	P
$\frac{3}{4}$	25	31	31	32	P
$1\frac{1}{2}$	41	22	27	34	P
$1\frac{3}{4}$	57	51	63	37	P
2	160	160	87	200	P
$2\frac{1}{4}$	510	520	—	—	—
$2\frac{1}{2}$	—	—	1200	—	—
$2\frac{3}{4}$	2.9	1.2	0.8	0.8	N

TABLE III — THE CHARACTERISTICS OF SOME HIGH VOLTAGE RECTIFIERS PROCESSED FROM GOLD DIFFUSED AND ZONE REFINED SILICON

Units ¹	E_R (volts) ²			E_F (volts) ³			τ (μ sec) ⁴	Silicon-Type
	10 μ a	100 μ a	1 ma	10 ma	100 ma	1A		
Me-512	30	120	500	0.8	1	1.3	1.6	Gold diffusion $\rho \sim 16,000 \Omega \text{ cm}$
513	22	200	600	1.0	1.2	1.5	0.6	
514	300	600	1000	0.8	1	1.5	2.1	
515	300	500	1000	0.8	1	1.4	1.2	
Me-375	1200	1500		2.5	3.5		<1	Floating zone refined $\rho \sim 6,000 \Omega \text{ cm}$
376	70	300	800	2.5	4.0		<1	
377	16	120	700	3.5	7			
378	320	400	800	2.5	3.5			

¹ These units have an area $\sim 10^{-3} \text{ cm}^2$.

² This is the reverse voltage at which these units pass the indicated current.

³ This is the forward bias at which the units pass the indicated current.

⁴ This is the lifetime measured at 30 milliamperes forward current by the pulse injected technique. The lifetime did not seem very sensitive to small variations in injected current.

the floating zone apparatus.¹³ This technique removes impurities from molten silicon by treatment with hydrogen containing water vapor. The material obtained from this process has an impurity level in the range of 10^{12} to 5×10^{13} acceptors/cm³ (2,000 to 16,000 $\Omega \text{ cm}$ P type).

Table III gives the characteristics of some of the better diodes made from such floating zone silicon. The reverse currents are larger than that predicted by (2.10). The lifetime at high injection is in the order of 1 μ sec.

N-type silicon with a resistivity range of 10 to 30 $\Omega \text{ cm}$ was diffused with gold at 1,200° for sixteen hours. With this diffusion program the gold is uniformly distributed in the material.¹⁴ The resistivity after gold diffusion was in the range of 2,000 to 15,000 $\Omega \text{ cm}$. The characteristics of several devices processed from this material are given in Table III. This technique has many attractive features; however, additional work was not done because the lifetime in the diffused material was consistently lower than that required for conductivity modulation.

One successful purification technique is horizontal zone refining of silicon in a quartz boat. With the number of passes used, the background acceptor concentration is observed to be in the order of 5×10^{13} to 10^{14} (100–1,000 $\Omega \text{ cm}$ P type). Most of the devices reported in this paper are fabricated from this material.

Capacity data for devices fabricated from various types of high resistivity silicon is shown in Fig. 11. The plot shows that the high resistivi-

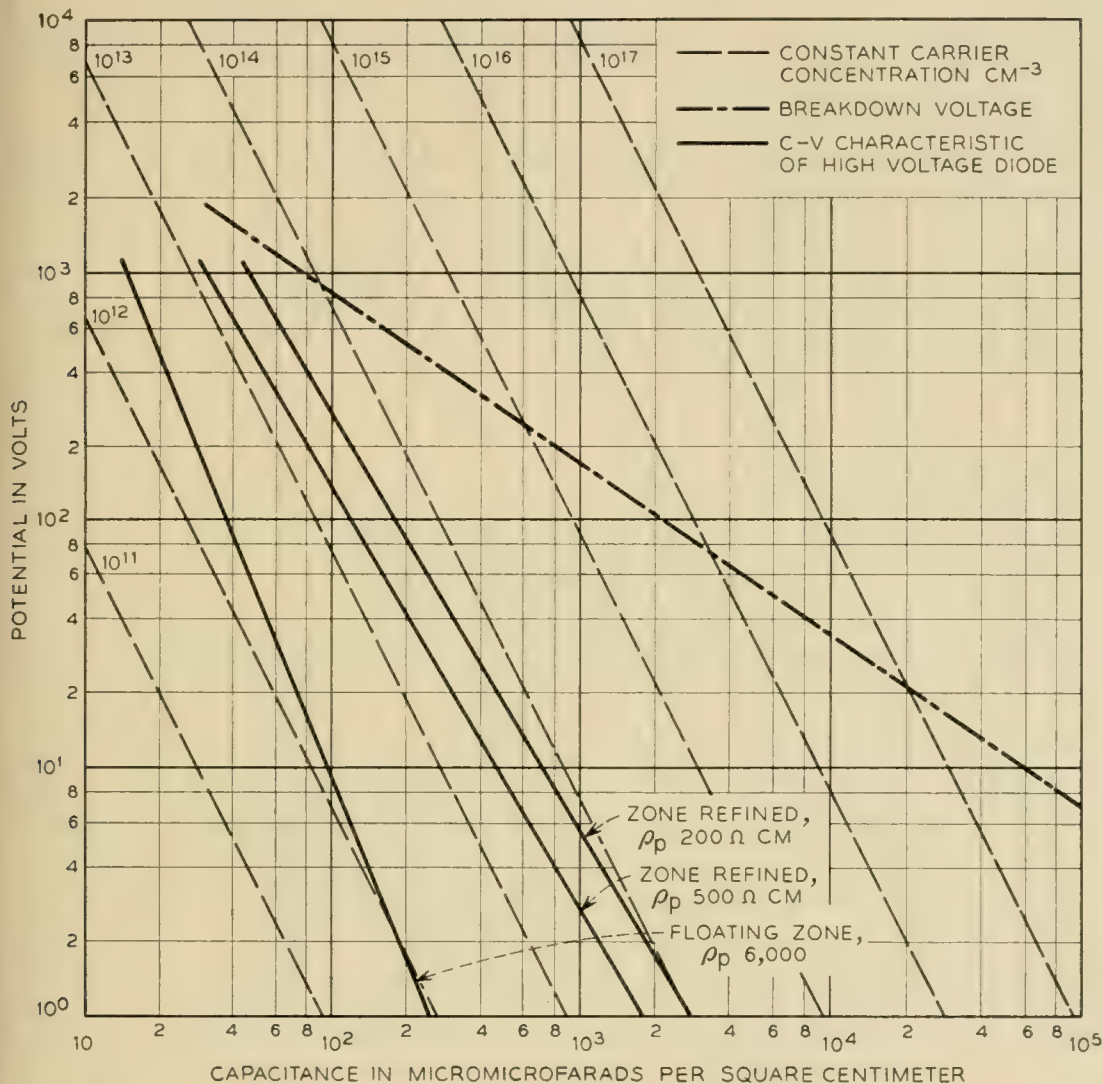


Fig. 11 — Capacity/cm² for high voltage rectifiers processed from various types of high resistivity material.

ties measured by the four point probe before diffusion are indicative of the impurity level after processing. The water vapor floating zone refined material has an impurity level of 10^{12} acceptors/cm³; the other material is in the range of 10^{13} to 10^{14} . The breakdown voltage line was calculated from the data in Fig. 5.

5.2 Diffusion

In this section some of the practical difficulties observed in utilizing the diffusion technique will be considered. In the fabrication of transistors close geometry control is necessary in order to obtain the desired device characteristic. It has been shown¹ that in conductivity modulated rectifiers the only geometry requirement is that the width of the center

region be less than the diffusion length of the minority carrier. High surface concentration of diffusant is desirable since this facilitates the contact problem. This suggests that the diffusion system can be much less involved than that required for diffused transistors.¹⁵ Some of the data presented in this section will show that the open tube diffusion technique¹⁶ can lead to variations in diffusion parameters.

The diffusion of impurities into silicon is complicated by variations in the boundary conditions at the surface. Frosch¹⁵ has shown that surface concentrations can be varied over six decades.

5.2.1 Device Diffusion Theory

Several important impurity distributions have been considered¹⁷. Two distributions are important in the open tube process:

1. *Error Function Complement, ERFC, Distribution or Infinite Diffusant Source*. If the diffusant is deposited on the silicon and serves as an infinite source, the added impurities will have an *erfc* distribution. For one diffusant and a fixed diffusion program this distribution will result in the deepest penetrations and smallest sheet resistances of all possible distributions. The sheet resistance is a measure of the total number of added impurities. The data presented later indicates that the added impurities frequently have an *erfc* distribution.

2. *Gaussian and Modified Gaussian Distribution*. A number of impurity atoms enters the solid, and a surface barrier builds up with time which prevents additional atoms from entering.¹⁷ Initially, the diffusant is assumed to be present in an infinitely thin layer at the surface with diffusion into or out of the material possible. In the range of silicon doping levels and surface concentrations used, a Gaussian, modified Gaussian or *erfc* distribution for a given diffusion program lead to approximately equal junction depths.

The sheet resistance and the diffusion depth have been related¹⁸ to the surface concentration for an *erfc* distribution. If the distribution is Gaussian instead of *erfc*, then for the same value of sheet resistance and diffusion depth the surface concentration should be reduced by one-third. Since the sheet resistance is related to the total number of impurities through a mobility term, quantitative interpretation of the data for any case other than *erfc* or Gaussian distributions would be difficult.

5.2.2 Experimental Results

The sheet resistance was measured by the four-point probe method, and the diffusion depths, by angle lapping and staining.¹⁹ Surface con-

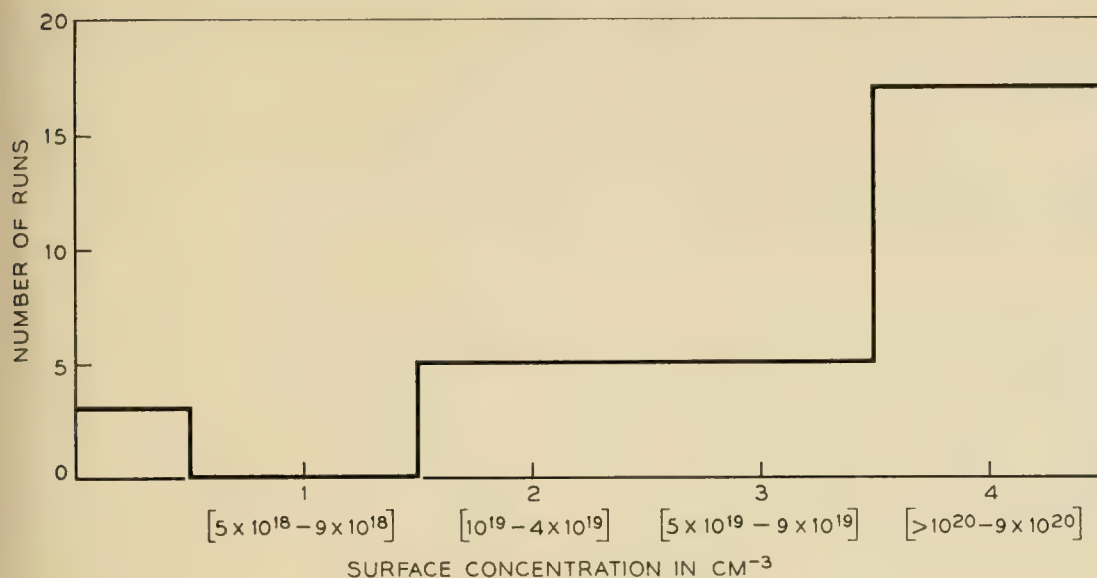


Fig. 12 — Distribution of surface concentration of 28 P_2O_5 diffusions by the open tube deposition technique.

centrations were calculated assuming an *erfc* distribution.¹⁸ All the diffusions are on lapped silicon surfaces in the temperature range of 1,200 to 1,300° C.

Fig. 12 shows the distribution of surface concentration of 28 P_2O_5 diffusions by the open tube process.¹⁶ The surface concentrations vary from 10^{19} to 5×10^{20} atoms/cm³. These values are about a decade lower than the closed tube values of surface concentrations reported by Fuller.¹⁹

The measured diffusion depths were in the order of 2×10^{-3} to 5×10^{-3} cm. Fig. 13 shows the distribution of diffusion depths normalized with the calculated diffusion depth as unity. The diffusion depths were calculated from the measured surface concentration assuming an *erfc*¹⁹ distribution.

The observed variation in diffusion depth is difficult to explain. Some of the possibilities which have been considered are:

1. The diffusion temperature from lot to lot would have to be from 0 to 50 degrees below the expected value to explain the variations. Discrepancies this large have not been observed.

2. One impurity distribution which may explain some of the results is a modified Gaussian with considerable out diffusion. There are some runs with high sheet resistance and diffusion depths which are consistent with this picture. Generally the sheet resistances are so small that there could not be much out diffusion.

3. Some workers have suggested the possibility of the diffusion constant being a function of the surface concentration. Fig. 13 does not

indicate any correlation between surface concentration and diffusion constant.

The variations in diffusion process control have not been observed to effect the production of rectifiers. If better geometry control is necessary, more sophisticated diffusion techniques are required.

VI PULSE PROPERTIES AND RELIABILITY

Important considerations in all diode applications are the pulse properties and reliability in operation. In this section some problems which are associated with avalanche breakdown are described and the results related to recent work on surface and body breakdown.

6.1 Theory

Several workers²⁰ have considered the possibility of a negative resistance in the avalanche region for reverse biased junctions in which one side is either intrinsic or so weakly doped that the space charge of the carriers cannot be neglected. A negative resistance might be observed at very high current densities in an $N^+\pi$ junction.

One possible source of a negative resistance would be a large temperature rise due to current concentration at a few points instead of a uni-

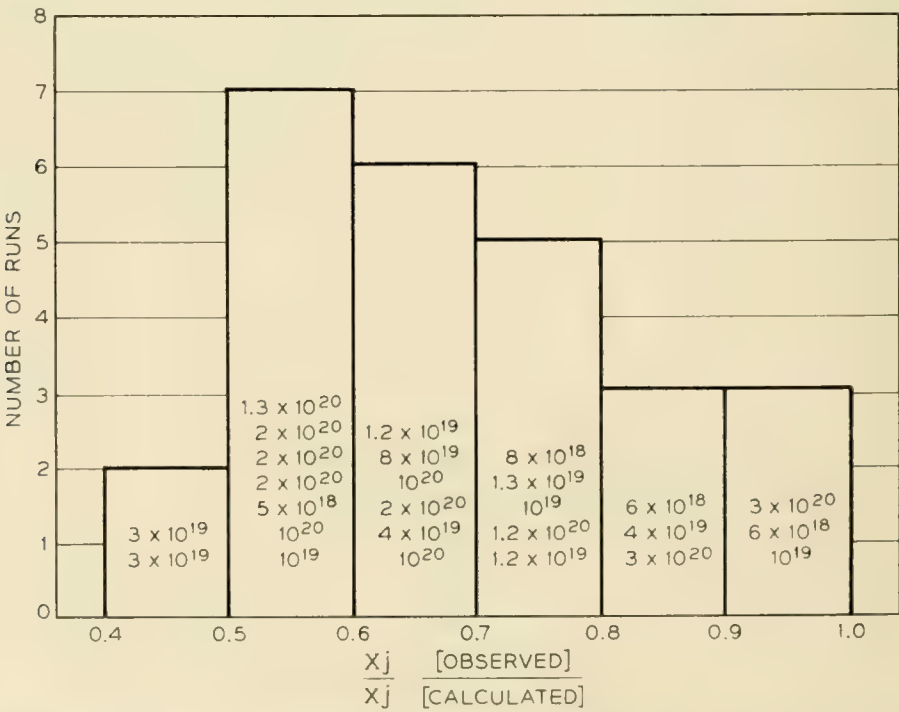


Fig. 13 — Distribution of diffusion depths for diffusion by the open tube deposition technique.

form flow through the junctions. This case is of particular significance in high voltage rectifiers where small reverse currents result in relatively large power. It has been pointed out in Sections 3.2 and 5.1 that body avalanche breakdown is frequently not observed in these devices.

Avalanche breakdown current in silicon⁸ is carried by discrete pulses of about $50 \mu a$ at their onset and increasing with increasing current to about $100 \mu a$. Approximate calculations²¹ show that the ionizing regions of these microplasmas are about $500 \text{ \AA}$ in extent, have a current density $\approx 2 \times 10^6 \text{ amp/cm}^2$, and have a net space-charge density $\approx 10^{18}/\text{cm}^3$. These pulses for junctions with E_{max} less than 500 kv/cm appear to be independent of junction width and built-in space-charge. Rose considers the statistical problem associated with a large number of pulses and presents a picture which is consistent with most of the experimental data. He calculates the temperature rise, assuming the avalanche power is 1×10^{-2} watts and is dissipated uniformly in a sphere. The maximum temperature rise for a cluster of two or three pulses is in the order of 25°C . For the picture Rose presents, the temperature rise due to the microplasma should be relatively insensitive to the breakdown voltage. Thermal collapse of rectification, i.e., increase of temperature until the silicon is intrinsic, will probably not occur in the region of avalanche multiplication. Two important conclusions can be obtained:

1. Avalanche breakdown should occur as a random process with a uniform probability over the junction. Large temperature rises due to a breakdown of microplasma will probably not occur since the resulting temperature rise would cause the breakdown voltage in that spot to increase. The power is dissipated throughout the path of the current pulse in the space-charge region.

2. A thermal effect in silicon due to heating by the small plasma has a very short time constant of the order of 10^{-10} seconds.²¹ It is not possible to separate a thermal effect of this type by reducing the pulse width. The heating and cooling time is short compared to the pulse time in these experiments.

The pulse properties of a junction would be quite different if the breakdown occurred at one spot instead of many spots distributed over the junction. Breakdown at a single spot on the surface has been observed.²²

6.2 Experimental Results

Many rectifiers were given a voltage pulse which carried them into breakdown. There was a wide distribution of V-I characteristics. Many diodes did not show a negative resistance up to the maximum instan-

taneous power the pulser could deliver, 5kw. These diodes are not considered in the subsequent analysis.

The diodes were subjected to 50 μ sec triangular voltage pulses which would send them into breakdown. Variations in pulse conditions did not effect the I-V characteristic until large pulses destructively damaged the unit.

Fig. 14 is a sketch of a typical V-I characteristic and Fig. 15, shows the voltage-versus-time characteristic for a diode with a negative resistance. The V-I curve can be broken into four regions:

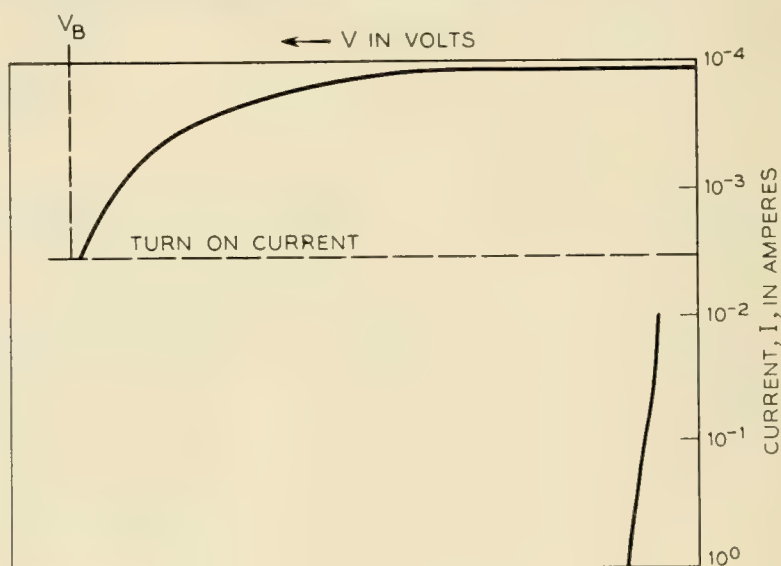


Fig. 14 — A typical V-I characteristic for a diode in which a negative resistance is observed.

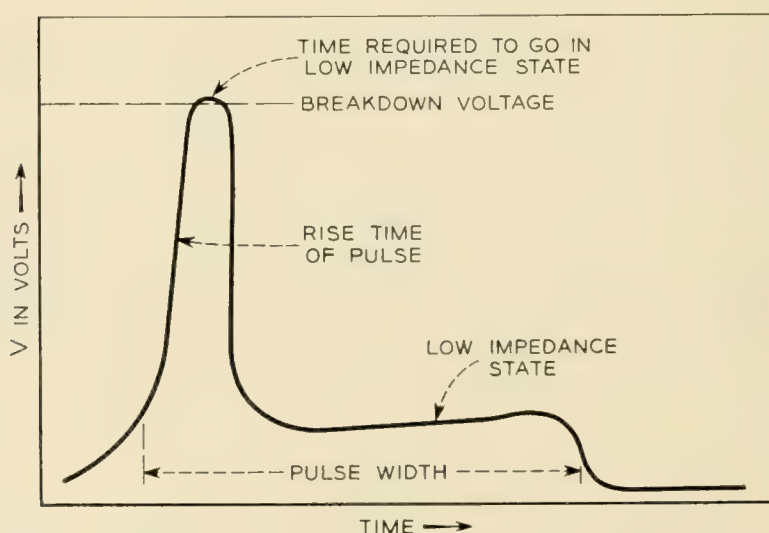


Fig. 15 — A typical V-T characteristic for a diode in which a negative resistance is observed.

1. A high impedance state before the breakdown voltage is reached.
2. A current required to turn on the negative resistance; this current varies from 10^{-3} to 1 amp.
3. The transition to a low impedance state.
4. The low impedance region in which the current is probably limited by the circuit impedance.

The V-T curve can be broken in four regions:

1. The time it takes the pulse to reach the breakdown voltage.
2. The time the diode can maintain the breakdown voltage less than $1 \mu\text{sec}$. This is beyond the resolution of the oscilloscope.
3. The time required to fall to the low voltage (low impedance) state, is less than $1 \mu\text{sec}$.
4. The remainder of the pulse in the low voltage state.

Fig. 16 is a plot showing the current and voltage required to turn on a negative resistance in several power rectifiers (area $\sim 10^{-2} \text{ cm}^2$). To

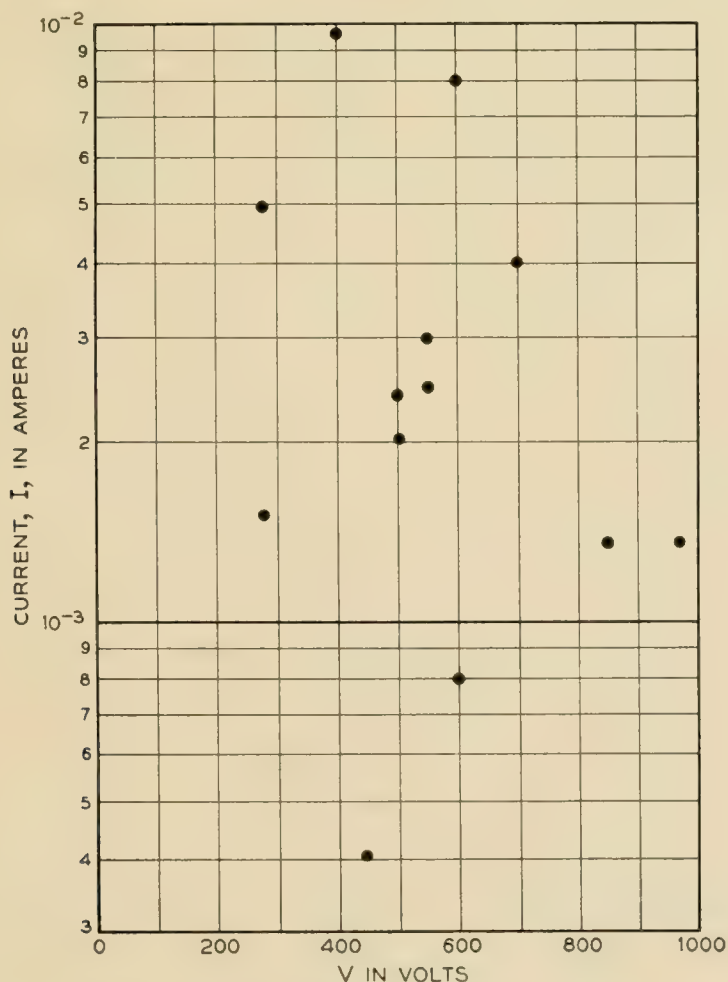


Fig. 16 — Current and voltage required to turn on a negative resistance in several power rectifiers ($A \sim 10^{-2} \text{ cm}^2$).

consider the spread of breakdown voltage, the data was normalized to the instantaneous power required to turn on a negative resistance. This turn-on power was the turn-on power multiplied by the voltage.

Fig. 17 is a plot of the distributions of turn-on power for the rectifiers which had a negative resistance plotted as a log normal distribution on probability paper. The median value for the turn-on power is 1.2 watts. Eighty per cent of these diodes went into a negative resistance condition at powers between 0.1 and 10 watts. Many diodes could dissipate several kilowatts with no negative resistance. These were not included.

Experiments show that devices which show surface breakdown will collapse at power levels which are orders of magnitude below that observed for devices in which body breakdown is observed.

The picture is more cloudy with smaller area rectifiers (area $\sim 10^{-3}$ cm²). In these devices it was not possible to predict the pulse properties

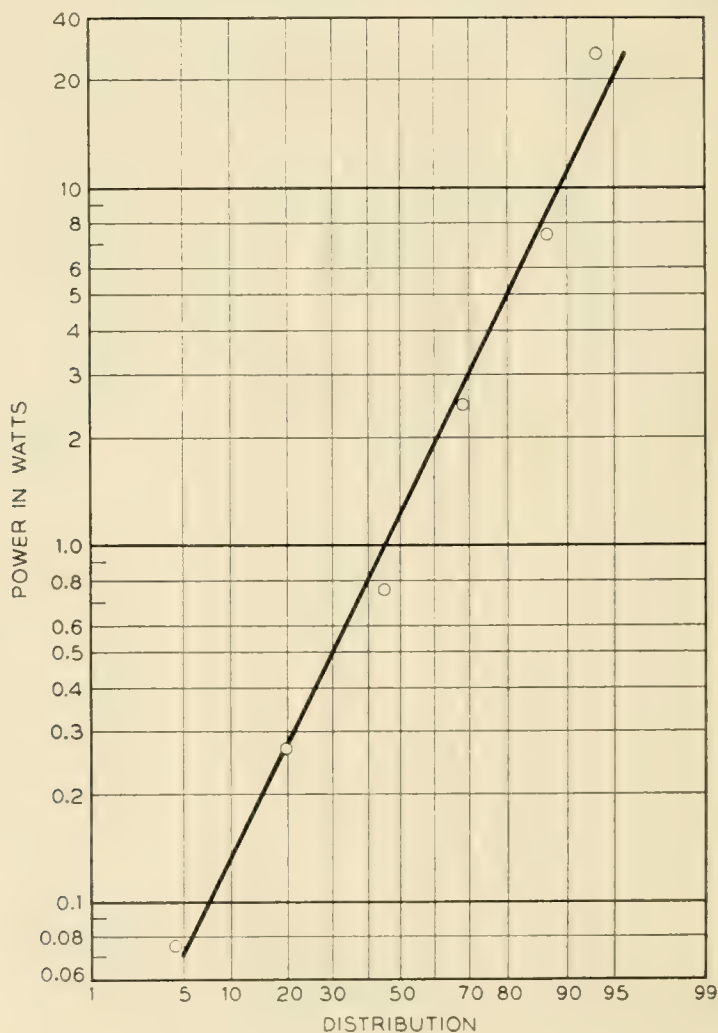


Fig. 17 — The distribution of turn-on power for rectifiers ($A \sim 10^{-2}$ cm²) in which a negative resistance is observed plotted as a log normal distribution on probability paper.

of the device from the reverse I-V characteristic. This may be attributed to the decrease in power capabilities of the body breakdown process in the smaller devices. This also suggests that the smaller devices have a less severe surface problem.

The distribution of turn-on power for a few hundred small area rectifiers ($A \sim 10^{-3} \text{ cm}^2$) is shown in Fig. 18. The median of the distribution occurs at 40 watts. Eighty percent of the units will show a negative resistance when pulsed at power levels between 3 and 500 watts.

VII CONCLUSION

High voltage rectifiers have been fabricated using several sources of high resistivity material employing an uncomplicated diffusion process.

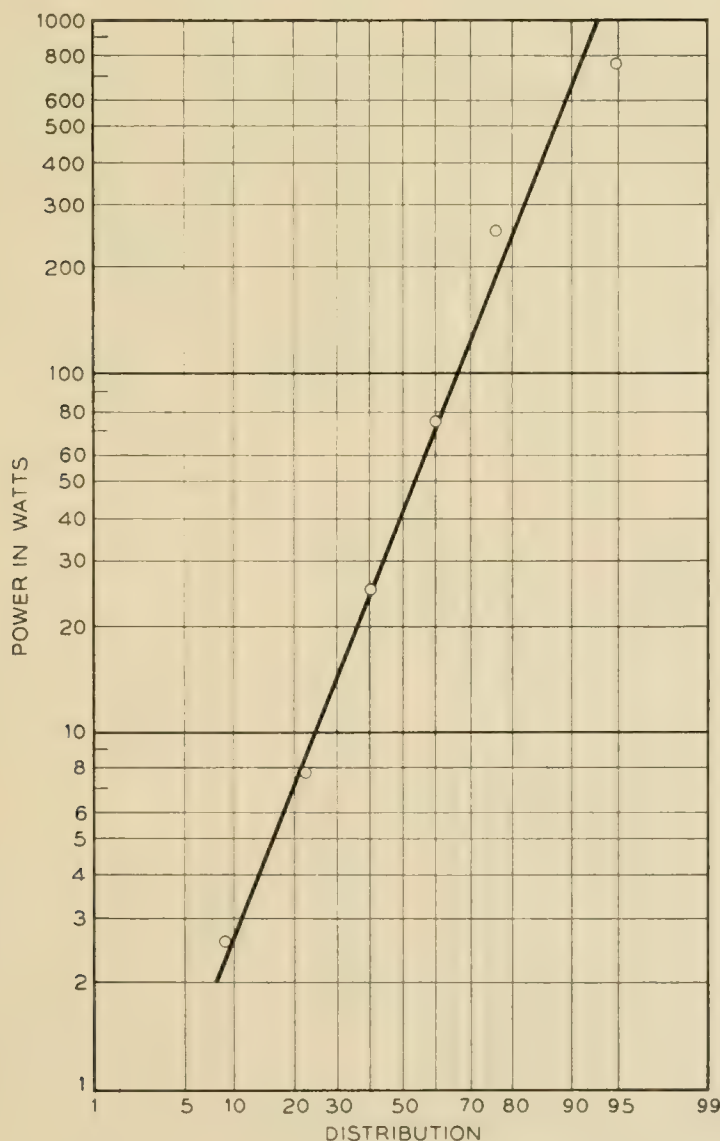


Fig. 18 — The distribution of turn-on power for small area rectifiers ($A \sim 10^{-3} \text{ cm}^2$) plotted as a log normal distribution on probability paper.

The most consistent results were obtained using horizontal zone refined silicon. The open tube diffusion technique has sufficient control to satisfy the fabrication requirements.

The magnitudes, voltage and temperature dependences of both the forward and reverse currents of silicon rectifiers can be explained by including a recombination level near the middle of the forbidden energy gap. Design equations for the forward and reverse characteristic of a diode are presented for several important cases. The breakdown voltage of the high voltage devices was shown to be a function of the width of the high resistivity region.

One unsolved problem is the surface limitation of breakdown voltage and reverse currents. This has been observed to decrease the breakdown voltages and increase the reverse currents to undesirable levels.

ACKNOWLEDGEMENT

The authors wish to thank their colleagues for many helpful discussions. Thanks are due Dr. R. N. Noyce of the Shockley Semiconductor Laboratory for making his paper available before publication. The work on zone refined and compensated silicon was done by S. J. Silverman. The floating zone material was supplied by H. E. Bridgers. Much of the experimental work was done by A. R. Tretola, T. J. Vasko and F. R. Lutchko. G. J. Levenbach assisted with the statistical aspects.

REFERENCES

1. M. B. Prince, B.S.T.J., **35**, p. 661, 1956. R. N. Hall, Proc. I.R.E., **40**, p. 1512, 1952.
2. W. Shockley, and W. T. Read, Jr., Phys. Rev., **87**, p. 835, 1952.
3. R. N. Hall, Proc. I.R.E., **40**, p. 1512, 1952. J. S. Saby, Proc. Rugby Conference on Semiconductors, 1956.
4. W. Shockley, B.S.T.J., **28**, p. 435, 1949.
5. K. G. McKay, Phys. Rev., **94**, p. 877, 1954.
6. E. M. Pell, J. Appl. Phys., **26**, p. 658, 1955. E. M. Pell, and G. M. Roe, J. Appl. Phys., **27**, p. 768, 1956.
7. B. Ross, and J. R. Madigan, Bull. A.P.S., **2**, p. 65, 1957.
8. K. G. McKay, Phys. Rev., **94**, p. 877, 1954.
9. S. L. Miller, Phys. Rev., **99**, p. 1234, 1955.
10. S. L. Miller, Phys. Rev., **105**, p. 1246, 1957.
11. P. A. Wolff, Phys. Rev., **95**, p. 1415, 1954.
12. C. T. Sah, R. N. Noyce, and W. Shockley, "Carrier Generation and Recombination in p - n Junctions and p - n Junction Characteristics", to be published in the Proc. I.R.E.
13. H. C. Theuerer, J. Metals, p. 1316-1319, Oct. 1956.
14. J. A. Burton, Physica, **20**, p. 845-854, 1954.
15. C. J. Frosch and L. Derick, J. Elec. Chem. Soc., to be published.
16. K. D. Smith, P.G.E.D. Conference of the I.R.E., Washington, 1956.
17. F. M. Smits and R. C. Miller, Phys. Rev., **104**, p. 1242-45, 1956.
18. G. Backenstoss, to be published.
19. C. S. Fuller and J. A. Ditzemberger, J. Appl. Phys., **27**, p. 544-53, 1956.
20. W. T. Read, Jr., B.S.T.J., **35**, p. 1239, 1956.
21. D. J. Rose, Phys. Rev., **105**, p. 413, 1957.
22. C. G. B. Garrett and W. H. Brattain, J. Appl. Phys., **27**, p. 299-306, 1956.

Coincidences in Poisson Patterns

By E. N. GILBERT and H. O. POLLAK

(Manuscript received August 3, 1956)

A number of practical problems, including questions about reliability of Geiger counters and short-circuits in electric cables, reduce to the mathematical problem of coincidences in Poisson patterns. This paper presents the probability of no coincidences as well as asymptotic formulas and simple bounds for that probability under a variety of circumstances. The probability of exactly N coincidences is also found in some cases.

INTRODUCTION

A number of practical problems are questions about what we call "coincidences" in Poisson patterns. In d -dimensional space, a Poisson pattern of density λ is a random array of points such that each infinitesimal volume element, dV , has probability λdV of containing a point, and such that the numbers of points in disjoint regions are independent random variables. Then a volume, V , has probability

$$\frac{(\lambda V)^k}{k!} e^{-\lambda V}$$

of containing exactly k points. A coincidence, in our usage of the word, is defined as follows: We imagine a certain fixed distance δ to be given in advance; two points are then said to be *coincident* if they lie within distance δ of one another.

Examples

The best-known case of a coincidence problem concerns Geiger counters. In the simplest mathematical model, there is a short dead-time δ after each count during which other particles can pass through the counter without registering a count. In our present terminology, a count is missed whenever two particles traverse the counter with coincident times of arrival. The same problem is encountered with telephone call registers.

Another example arises in the manufacture of electric cable. Each wire in a cable is covered with an insulation which contains occasional flaws. When the cable is assembled it will fail a short circuit test if it contains a pair of wires such that a flaw on one wire lies within some distance δ of a flaw on the other wire. In a similar way, coincident flaws in the insulation of the wire from which a coil is wound can lead to failure of the coil.

There are also some problems in the development of certain military systems which lead to the consideration of coincidences in Poisson patterns.

Outline of Work

Our primary aim is to study the probability of no coincidences under various circumstances. In Part I, we examine coincidences of two different Poisson patterns, of densities λ and μ respectively, on a line of length L . Here we do not count two points of the *same* pattern within a distance δ as giving a coincidence. A set of integral equations yields the probability of no coincidences as well as an asymptotic formula and upper and lower bounds.

In Part II, we study the probability, $F_0(L)$, of no coincidences for a single one-dimensional Poisson pattern of density λ . These results may also be interpreted as the distribution function for the minimum distance between pairs of points of a Poisson pattern. Sample formulas are the asymptotic formula (for large L)

$$F_0(L) \approx \frac{\lambda}{(\lambda - a)[1 + \delta(\lambda - a)]} e^{-aL},$$

and the bounds (valid for all L)

$$\left(1 - \frac{a}{\lambda}\right) e^{-aL} e^{-(a-\lambda)\delta} \leq F_0(L) \leq e^{-aL} e^{-(a-\lambda)\delta},$$

where $s = -a$ is the largest real root of

$$s + \lambda = \lambda e^{-(s+\lambda)\delta}.$$

The problem of n Poisson patterns, all of the same density λ , is examined in Part III. Coincidences are now counted between points of any two distinct patterns.

The one-dimensional problems of Parts I–III succumb readily to analytic techniques. We can find exact expressions for the probabilities of no coincidences in Parts I–III. Two entirely different methods of deriving

exact results are available and are illustrated in Parts II and III. Unfortunately, the exact formulas, although they are finite sums, contain a number of terms which grows with L . Much of our effort has been directed toward finding good, easily computed bounds and asymptotic formulas.

The probabilities of having exactly N coincidences are also obtainable but they have more complicated formulas. A detailed derivation is given only in Part II.

In Part IV, we consider the probability of no coincidence in higher dimensional problems. The methods of Parts I–III fail in higher dimensions, but we are still able to derive some bounds. An exact formula is derived for the probability of no coincidences within a single two-dimensional Poisson pattern in a rectangle with sides $\leq 2\delta$. We also give particular attention to coincidences in a three-dimensional cylinder.

Part V contains numerical results.

Reduction of the Examples to the Theory

We now wish to see how answers bearing on the practical problems previously listed may be found from this study.

The literature on Geiger counters (see bibliography in Feller³) is concerned with statistics of the number of counts registered in a given long time, t . The basic problem is to test the hypothesis that the particles arrive in a Poisson sequence. To this problem, then, are relevant the formulas for the probability of N coincidences in one pattern given in Part II, and the bounds and asymptotic results there derived.

The problem of coincident flaws in an electric cable is three-dimensional, and we have various approaches leading to the probability of no coincidences which are valid under different circumstances. If the cable contains only two wires (with possibly different flaw densities), then the problem reduces to the one-dimensional case of coincidences between two Poisson patterns treated in Part I. If the diameter of the cable is small with respect to δ , and if the density of flaws is the same on each of the n wires in the cable, we have the situation of n identical patterns treated in Part III. If, in addition, n is very large, we may ignore the fact that coincident flaws on a single wire do not cause short circuits, and think of coincidences within a single pattern (Part II). Without the assumption that the diameter of the cable is small with respect to δ , the problem is no longer reducible to a one-dimensional form. Section 4.4 is especially devoted to thick cable, and to producing a lower bound for the probability of no coincidences in this three-dimensional situation.

The literature on Poisson patterns in a line segment contains the fol-

lowing related papers. C. Domb¹ finds the distribution function for the total length of the set of points lying within distance δ of a pattern point. P. Eggleton and W. O. Kermack² and also L. Silberstein⁵ consider *aggregates*, which are sets of k pattern points all contained in an interval of length δ . In the special case $k = 2$, aggregates are our coincidences. These authors find the expected number of aggregates but not the probability of N aggregates.

I COINCIDENCES BETWEEN TWO PATTERNS

1.1 Integral Equation

Consider two Poisson patterns of points on the real line, the first with density λ (points per unit length) and the second with density μ . We want the probability $F(L)$ that in the segment from 0 to L there is no coincidence between a point of pattern No. 1 and a point of Pattern No. 2. $F(L)$ will be formulated in terms of the conditional probabilities

$P_1(L) = \text{Prob (no coincidence, given Pattern No. 1 has point at } L),$

$P_2(L) = \text{Prob (no coincidence, given Pattern No. 2 has point at } L).$

If $L \leq \delta$, $P_1(L)$ and $P_2(L)$ are the probabilities that patterns No. 2 and No. 1 are empty:

$$P_1(L) = e^{-\mu L}, \quad P_2(L) = e^{-\lambda L}, \quad \text{if } L \leq \delta. \quad (1-1)$$

If $L > \delta$ and Pattern No. 1 contains a point at L , there are two ways that no coincidences can occur. First, Pattern No. 2 may fail to have any points anywhere in the interval $[0, L]$. The probability of this event is $\exp - \mu L$. The second possibility is illustrated in Fig. 1 (using circles for points of Pattern No. 1 and crosses for points in Pattern No. 2). Pattern No. 2 has points in $(0, L)$; the one closest to L is at $y < L - \delta$. Since the interval (y, L) contains no points of Pattern No. 2, the probability of finding this closest point, y , in an interval, dy , is

$$\exp [-\mu(L - y)]\mu dy.$$

The interval $(y, y + \delta)$ must be free from points of Pattern No. 1 (prob-

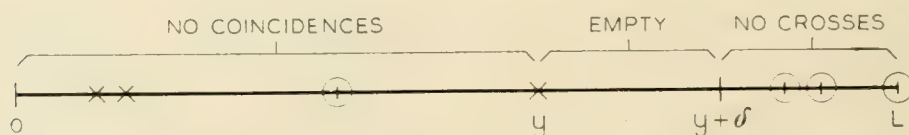


Fig. 1 — Patterns without coincidence.

ability $\exp(-\lambda\delta)$) and the interval $(0, y)$ must contain no coincidences (probability $P_2(y)$). One obtains finally

$$P_1(L) = e^{-\mu L} \left[1 + \mu e^{-\lambda\delta} \int_0^{L-\delta} e^{\mu y} P_2(y) dy \right], \quad (1-2)$$

and similarly

$$P_2(L) = e^{-\lambda L} \left[1 + \lambda e^{-\mu\delta} \int_0^{L-\delta} e^{\lambda y} P_1(y) dy \right]. \quad (1-3)$$

The solutions $P_1(L)$ and $P_2(L)$ are determined uniquely by (1-1), (1-2) and (1-3). For (1-1) determines them for $0 \leq L \leq \delta$ and the integrations indicated in (1-2) and (1-3) will provide the solutions in $0 \leq L \leq (n+1)\delta$ when they are known in $0 \leq L \leq n\delta$. $P_1(L)$ and $P_2(L)$ are piecewise analytic; the analytic form of the solution changes each time L passes an integer multiple of δ . These analytic expressions soon become complicated and are less useful than the bounds and approximations given later on.

To compute $F(L)$, consider the last place before L at which either Pattern No. 1 or No. 2 has a point. The probability that this last point lies between x and $x+dx$ and belongs to Pattern No. 1 is $\exp[-(\lambda+\mu)(L-x)]\lambda dx$ (Fig. 2). This term multiplied by $P_1(x)$ and integrated from 0 to L gives the probability of no coincidences if the last point is a circle. A similar integral gives the probability if the last point is a cross. Finally there is probability $\exp[-(\lambda+\mu)L]$ that neither pattern has a last point [i.e., $(0, L)$ empty]. Then

$$F(L) = e^{-(\lambda+\mu)L} \left[1 + \int_0^L e^{(\lambda+\mu)x} (\lambda P_1(x) + \mu P_2(x)) dx \right]. \quad (1-4)$$

1.2 Solution by Laplace Transforms

For $i = 1$ or 2 , let

$$p_i(s) = \int_0^\infty P_i(L) e^{-sL} dL. \quad (1-5)$$

Replacing $P_1(L)$ in (1-5) by (1-1) for $0 \leq L \leq \delta$, by (1-2) for $\delta \leq L$,

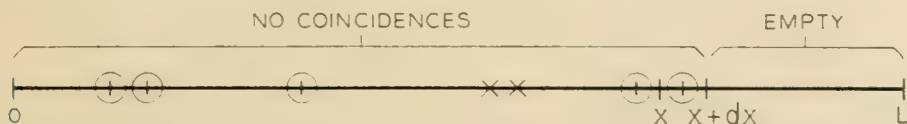


Fig. 2 — Patterns without coincidence.

and interchanging the order of integration of a double integral,

$$(s + \mu)p_1(s) = 1 + \mu e^{-(\lambda + \mu + s)\delta} p_2(s).$$

Similarly,

$$(s + \lambda)p_2(s) = 1 + \lambda e^{-(\lambda + \mu + s)\delta} p_1(s),$$

so that

$$p_1(s) = \frac{s + \lambda + \mu e^{-(\lambda + \mu + s)\delta}}{(s + \lambda)(s + \mu) - \lambda \mu e^{-2(\lambda + \mu + s)\delta}}, \quad (1-6)$$

and

$$p_2(s) = \frac{s + \mu + \lambda e^{-(\lambda + \mu + s)\delta}}{(s + \lambda)(s + \mu) - \lambda \mu e^{-2(\lambda + \mu + s)\delta}}. \quad (1-7)$$

Likewise, using (1-4), the Laplace transform $f(s)$ of $F(L)$ is

$$f(s) = \frac{1 + \lambda p_1(s) + \mu p_2(s)}{\lambda + \mu + s}.$$

As one might expect from the piecewise analytic character of $P_1(L)$ and $P_2(L)$ there is no convenient way of transforming $f(s)$ back to $F(L)$. By evaluating residues of $f(s) \exp(sL)$ at the poles of $f(s)$ one might express $F(L)$ as an infinite series of exponential terms. The most slowly damped term in this series can be expected to approximate $F(L)$ when L is large. The poles of $f(s)$ are at the zeros of the denominator $D(s)$ of $p_1(s)$ and $p_2(s)$:

$$D(s) = (s + \lambda)(s + \mu) - \lambda \mu e^{-2(\lambda + \mu + s)\delta}. \quad (1-9)$$

Since $D(x) > 0$ for $x \geq 0$ and both $D(-\lambda)$ and $D(-\mu)$ are negative, it follows that $D(s)$ has a real zero $s = -a$ with $a < \text{Min}(\lambda, \mu)$.

The zero $s = -a$ of $D(s)$ is the one with the largest real part. For, letting $s = x + iy$, we have in the half plane $x \geq -a$

$$\begin{aligned} |(s + \lambda)(s + \mu) - \lambda \mu e^{-2(\lambda + \mu + s)\delta}| &= |s + \lambda| \cdot |s + \mu| - \lambda \mu e^{-2(\lambda + \mu + x)\delta} \\ &\geq (x + \lambda) \cdot (x + \mu) - \lambda \mu e^{-2(\lambda + \mu + x)\delta} \geq 0. \end{aligned}$$

Also, if $y \neq 0$ the $\geq$ sign in the above proof can be replaced by $>$ and one concludes that all other zeros of $D(s) = 0$ satisfy

$$\text{Re } s < -b$$

for some $b > a$ (note that the left hand side of the preceding inequality does not approach 0 as y approaches $\pm \infty$).

The pole of $f(s)$ at $s = -a$ contributes to $F(L)$ a dominant term

$$F(L) \approx \frac{\lambda^2 + \mu^2 - (\lambda + \mu)a + 2\lambda\mu e^{-(\lambda+\mu-a)\delta}}{(\lambda + \mu - a)[\lambda + \mu - 2a + 2\delta(\lambda - a)(\mu - a)]} e^{-aL}. \quad (1-10)$$

In (1-10) the error is $O(\exp - bL)$ for large L .

When δ is small, we find $a = 2\lambda\mu\delta + O(\delta^2)$ and (1-10) becomes

$$F(L) \approx [1 + O(\delta^2)] \exp - [2\lambda\mu\delta + O(\delta^2)]L. \quad (1-11)$$

It is interesting to note that a simple heuristic argument also leads to a formula like (1-11). When δ is small and L is large, one expects that the intervals of length 2δ which contain points of Pattern No. 1 at their centers will comprise a total length near $(\lambda L)(2\delta)$ of the line segment $(0, L)$. The probability that a set of length $2\lambda L\delta$ shall be free of points of Pattern No. 2 is $\exp - 2\lambda\mu\delta L$.

1.3 Bounds

In this section we derive some relatively simple expressions which are good upper and lower bounds on $F(L)$. Both bounds have the same functional form:

$$K(A, B; L) = \frac{\lambda A + \mu B}{\lambda + \mu - a} e^{-aL} + \left(1 - \frac{\lambda A + \mu B}{\lambda + \mu - a}\right) e^{-(\lambda+\mu)L}. \quad (1-12)$$

In (1-12), a is again the smallest real solution of $D(-a) = 0$. A and B are positive constants which are related by

$$\frac{A}{B} = \frac{\mu}{\mu - a} e^{-(\lambda+\mu-a)\delta} = \frac{\lambda - a}{\lambda} e^{(\lambda+\mu-a)\delta}. \quad (1-13)$$

$K(A, B; L)$ becomes an upper bound or a lower bound depending on additional restrictions which will be placed on A and B .

To get the lower bound, we restrict A and B by the inequalities

$$A < e^{(a-\mu)\delta}, \quad B < e^{(a-\lambda)\delta}, \quad (1-14)$$

and

$$A < \left(1 - \frac{a}{\lambda}\right) e^{\mu\delta}, \quad B < \left(1 - \frac{a}{\mu}\right) e^{\lambda\delta}. \quad (1-15)$$

We first prove that (1-13), (1-14), and (1-15) imply

$$P_1(L) > Ae^{-aL}, \quad P_2(L) > Be^{-aL}. \quad (1-16)$$

When $0 \leq L \leq \delta$, (1-16) holds because of (1-1), (1-14), and the inequalities $a < \lambda$, $a < \mu$. If (1-16) were not true for all L there would be a smallest value, say $L = X > \delta$, at which at least one of the inequalities (1-16) would become an equality. Suppose the inequality (1-16) on $P_1(X)$ fails. Using (1-16) for $L < X$, and (1-2),

$$\begin{aligned} P_1(X) &> e^{-\mu X} \left(1 + B\mu e^{-\lambda\delta} \frac{e^{(\mu-a)(X-\delta)} - 1}{\mu - a} \right) \\ &> Ae^{-aX} + \left(1 - B \frac{\mu e^{-\lambda\delta}}{\mu - a} \right) e^{-\mu X} \quad \text{by (1-13),} \\ &> Ae^{-aX} \quad \text{by (1-15).} \end{aligned}$$

This contradicts our assumption that (1-16) fails for $P_1(X)$. A similar proof shows (1-16) cannot fail for $P_2(X)$.

Having proved (1-16) we now substitute these bounds into (1-4) and integrate to get $F(L) > K(A, B; L)$.

To make (1-12) into an upper bound it is only necessary to replace (1-14) and (1-15) by

$$A > 1, \quad B > 1, \quad (1-17)$$

and

$$A > \left(1 - \frac{a}{\lambda} \right) e^{\mu\delta}, \quad B > \left(1 - \frac{a}{\mu} \right) e^{\lambda\delta}. \quad (1-18)$$

The proof that now $F(L) < K(A, B; L)$ proceeds exactly as before but with all the inequality signs reversed.

Both bounds are dominated by an exponential term $\exp - aL$, as is the asymptotically correct formula (1-10). In typical numerical cases the coefficients multiplying this term in the three formulas agree closely. A numerical case is given in Part V.

1.4 Probability of N Coincidences

The methods of Sections 1.1 and 1.2 can also be used to find the probability $F_N(L)$ that there be exactly N coincidences in the interval $(0, L)$. It might appear most natural to define N to be the number of pairs of points (x, z) , x from Pattern No. 1, z from Pattern No. 2, such that

$$|x - z| < \delta. \quad (i)$$

However, we add the additional requirement that x and z be "adjacent" points; i.e.

$$\text{the interval } (x, z) \text{ is empty.} \quad (ii)$$

For example, in Fig. 3, we would count $N = 6$ coincidences even though there are 18 pairs which satisfy (i). In cable problems it appears reasonable to count coincidences as above. If we assume that all flaws are equally bad, then a short circuit is likely to develop only across an adjacent coincidence; our N is the number of places on the cable at which a short circuit can form. Another interpretation is that the cable can be cut into exactly $N + 1$ pieces each of which contain no coincidences.

Let $P_{1,N}(L)$ be the conditional probability of having N coincidences in $(0, L)$ knowing that there is a point of Pattern No. 1 at L . The Laplace transform of $P_{1,N}(L)$ turns out to be the coefficient of t^N in a generating function of the form

$$p_1(t, s) = \frac{\lambda + s + \mu\Omega}{(\lambda + s)(\mu + s) - \lambda\mu\Omega^2},$$

where $\Omega = e^{-(\lambda+\mu+s)\delta}(1-t) + t$. Interchanging λ and μ one gets the generating function $p_2(t, s)$ for the Laplace transform of the probability $P_{2,N}(L)$ of N coincidences, given a point of Pattern No. 2 at L . Finally the Laplace transform of $F_N(L)$ is the coefficient of t^N in the generating function

$$f(t, s) = \frac{1 - e^{-(\lambda+\mu+s)\delta} + \lambda p_1(t, s) + \mu p_2(t, s)}{\lambda + \mu + s}.$$

Since $f(t, s)$ is a rational function of t , it is easy to find the coefficient of t^N . The poles of this function are again just zeros of $D(s)$. Now, however, the poles are higher order poles. For large L an asymptotic formula for $F_N(L)$ has the form $\exp - aL$ times a polynomial in L with degree depending on N .

For more details about this method we refer the reader to Part II where a similar, but less involved, calculation is carefully done.

II SELF-COINCIDENCES IN ONE POISSON PATTERN

2.1 Integral Equation

In this part we shall consider a single one-dimensional Poisson pattern with density λ and ask for the probability $F_N(L)$ that in the interval $(0, L)$ the pattern have exactly N coincidences. We count coincidences

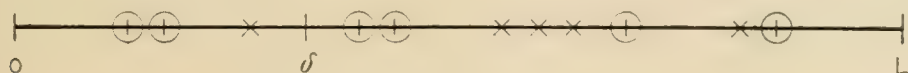


Fig. 3 — Patterns with six coincidences.

as in Section 1.4; a pair (x, z) of pattern points contributes one coincidence to the total number N only if both $|x - z| < \delta$ and the interval between x and z is empty.

Note that $F_0(L)$ is related to the distribution function for the minimum distance between the points of the pattern in $(0, L)$:

$$\text{Prob (min. dist.} \leq \delta) = 1 - F_0(L),$$

where it must be remembered that $F_0(L)$ is a function of δ .

As in Part I, we first define the conditional probabilities $P_N(L) = \text{Prob (exactly } N \text{ coincidences in } (0, L), \text{ given a point at } L)$. We then have the following equations:

$$\text{If } L \leq \delta, \quad P_N(L) = \frac{(\lambda L)^N}{N!} e^{-\lambda L}, \quad \text{all } N. \quad (2-1)$$

$$\text{If } L \leq \delta, \quad F_N(L) = \frac{(\lambda L)^{N+1}}{(N+1)!} e^{-\lambda L}, \quad N \geq 1. \quad (2-2)$$

If $L > \delta$, and $N \geq 1$, the probability of exactly N coincidences in $(0, L)$ equals the probability of N coincidences up to the last point of the pattern in the interval $(0, L)$ — and if there are to be any coincidences, there must be points of the pattern in $(0, L)$. Hence, if $L > \delta$, $N \geq 1$,

$$F_N(L) = \int_0^L P_N(L - y) e^{-\lambda y} \lambda dy. \quad (2-3)$$

If $N = 0$, the same argument applies, but there is also the possibility that there are no points at all of the pattern in $(0, L)$. Hence, if $L > \delta$,

$$F_0(L) = e^{-\lambda L} + \int_0^L P_0(L - y) e^{-\lambda y} \lambda dy. \quad (2-4)$$

Now let us consider the case where there is a point of the pattern at L . Then if the last point preceding L is between $L - \delta$ and L , this point and the point at L will create a coincidence; if there is no point within $(L - \delta, L)$, then all coincidences are within $(0, L - \delta)$. Hence, if $L > \delta$, and $N \geq 1$,

$$P_N(L) = \int_0^\delta P_{N-1}(L - y) \lambda e^{-\lambda y} dy + e^{-\lambda \delta} F_N(L - \delta). \quad (2-5)$$

For the case $N = 0$, we cannot allow a point in the interval $(L - \delta, L)$, and hence, if $L > \delta$,

$$P_0(L) = e^{-\lambda \delta} F_0(L - \delta). \quad (2-6)$$

2.2 Laplace Transform of $F_N(L)$

To analyze the system of equations which is given by relations (2-1) through (2-6), we introduce the generating functions

$$f(L, t) = \sum_{N=0}^{\infty} F_N(L) t^N,$$

and

$$p(L, t) = \sum_{N=0}^{\infty} P_N(L) t^N.$$

If $L > \delta$, we obtain from (2-3) and (2-4) the relation

$$e^{\lambda L} f(L, t) = 1 + \int_0^L p(w, t) e^{\lambda w} \lambda dw, \quad (2-7)$$

and from (2-5) and (2-6) the relation (again if $L > \delta$)

$$e^{\lambda L} p(L, t) = \lambda t \int_{L-\delta}^L p(w, t) e^{\lambda w} dw + e^{\lambda(L-\delta)} f(L - \delta, t). \quad (2-8)$$

If we differentiate (2-7) and (2-8) with respect to L , and then apply (2-7) differentiated to simplify the last terms of (2-8) differentiated, we obtain, still only for $L > \delta$,

$$f'(L, t) + \lambda f(L, t) = \lambda p(L, t), \quad (2-9)$$

$$p'(L, t) + \lambda(1 - t)p(L, t) = \lambda e^{-\lambda \delta} (1 - t)p(L - \delta, t). \quad (2-10)$$

It is easy to check from (2-1) and (2-2) that if $L \leq \delta$, then

$$p(L, t) = e^{-\lambda L(1-t)},$$

and

$$f(L, t) = e^{-\lambda L} \left(\frac{(e^{\lambda L t} - 1)}{t} + 1 \right),$$

and hence (2-9) is valid for *all* L , but the left side of (2-10) *vanishes* if $L \leq \delta$. Hence we may take Laplace transforms of (2-9) and (2-10). If we define

$$A(s, t) = \int_0^{\infty} f(L, t) e^{-Ls} dL,$$

and

$$B(s, t) = \int_0^{\infty} p(L, t) e^{-Ls} dL,$$

we obtain from (2-9), which we now know to be valid for all L ,

$$(\lambda + s)A(s, t) - 1 = \lambda B(s, t), \quad (2-11)$$

and from (2-10), by recalling that the left side vanishes for $L \leq \delta$,

$$sB(s, t) - 1 + \lambda(1 - t)B(s, t) = \lambda(1 - t)e^{-(s+\lambda)\delta}B(s, t). \quad (2-12)$$

Hence

$$B(s, t) = \frac{1}{s + \lambda(1 - t)[1 - e^{-(s+\lambda)\delta}]}, \quad (2-13)$$

and

$$A(s, t) = \frac{1}{\lambda + s} (1 + \lambda B(s, t)).$$

If we denote the Laplace transforms of $P_N(L)$ and $F_N(L)$ by $p_N(s)$ and $f_N(s)$ respectively, then

$$p_N(s) = \frac{\lambda^N [1 - e^{-(s+\lambda)\delta}]^N}{[s + \lambda - \lambda e^{-(s+\lambda)\delta}]^{N+1}}, \quad (2-14)$$

and

$$\begin{aligned} f_0(s) &= \frac{1}{\lambda + s} (\lambda p_0(s) + 1), \\ f_N(s) &= \frac{\lambda}{\lambda + s} p_N(s) \quad \text{for } N = 1, 2, \dots \end{aligned} \quad (2-15)$$

2.3 Exact Formula for $F_0(L)$

It is possible to solve (2-1) through (2-6) in piecewise analytic form by computing recursively from each interval of length δ to the next one. We shall obtain the piecewise analytic form for $F_0(L)$ by a direct derivation essentially due to E. C. Molina.⁴

Suppose k is the number of pattern points which fall into $(0, L)$. Let x_i denote the distance between the $i - 1^{\text{st}}$ point and the i^{th} point (x_1 is the distance from 0 to the first point) as shown in Fig. 4. The configura-

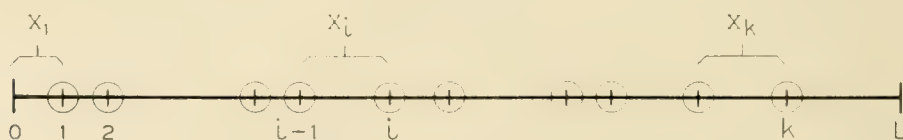


Fig. 4 — Definition of x_i .

tion of points $1, \dots, k$ on the line is represented by a single point $(x_1, \dots, x_k)$ in the polyhedron T in k -dimensional space defined by the inequalities

$$T: 0 \leq x_1, \dots, 0 \leq x_k, \quad x_1 + x_2 + \dots + x_k \leq L,$$

and the probability distribution of the point $(x_1, \dots, x_k)$ in T is uniform. The configurations with no coincidences lie in a smaller polyhedron T' consisting of all points of T for which $\delta \leq x_2, \dots, \delta \leq x_k$. Given k , the conditional probability that there be no coincidences is the ratio of two k -dimensional volumes $\text{Vol}(T')/\text{Vol}(T)$.

$$\text{Vol}(T') = 0 \quad \text{if} \quad L \leq (k-1)\delta.$$

For larger values of L let $y_1 = x_1, y_2 = x_2 - \delta, y_3 = x_3 - \delta, \dots, y_k = x_k - \delta$. Then T' becomes a polyhedron of the form

$$T'': 0 \leq y_1, 0 \leq y_2, \dots, 0 \leq y_k, \quad y_1 + y_2 + \dots + y_k \leq L - (k-1)\delta.$$

Since the transformation from x 's to y 's has determinant equal to one, T'' has the same volume as T' . However, T'' is now seen to be similar to T but with sides of length $L - (k-1)\delta$ instead of L . The volume ratio sought must be

$$\left(\frac{L - (k-1)\delta}{L} \right)^k.$$

Since k has the Poisson distribution with mean λL we obtain finally

$$F_0(L) = e^{-\lambda L} \sum_{k=0}^{1+[L/\delta]} \frac{(\lambda L)^k}{k!} \left(1 - \frac{(k-1)\delta}{L} \right)^k.$$

The piecewise-analytic character of $F_0(L)$ is evident; increasing L by an amount δ increases the upper limit on the sum by one and thereby adds a new term to the analytic expression for $F(L)$.

2.4 Asymptotic Formula for $F_N(L)$

Similar exact formulas could be found for all the $F_N(L)$, but they are both complicated and inconvenient for computing if L/δ becomes large. It is thus natural to aim for asymptotic results and for bounds connected with them.

The Laplace transform of $F_N(L)$ is given through (2-14) and (2-15) above. The pole of $f_N(s)$ with largest real part is a pole of order $N+1$

at a real negative point

$$s = -a > -\lambda.$$

For large L , the asymptotic behavior is given by

$$F_N(L) \approx \frac{\lambda e^{-aL}}{(\lambda - a)[1 + \delta(\lambda - a)]N!} \left[\frac{aL}{1 + \delta(\lambda - a)} \right]^N,$$

where the error term is $O(L^{N-1} e^{-aL})$ if $N \geq 1$. Such a formula, then, is a good approximation for fixed N as L increases; for fixed L , however, it will fail to be good for sufficiently large N .

If $N = 0$, the asymptotic form is

$$F_0(L) \approx \frac{\lambda}{(\lambda - a)[1 + \delta(\lambda - a)]} e^{-aL},$$

but the error term now decreases at a more rapid rate, as may be seen by including the contributions of some of the complex poles of $f_0(s)$. To find these poles, set

$$s + \lambda = \lambda e^{-(s+\lambda)\delta}.$$

If

$$s = -\lambda + r \exp(i\theta),$$

one obtains the simultaneous real system

$$2\pi m - \theta = \delta r \sin \theta \quad (m \text{ integer}),$$

$$\log(r/\lambda) = -\delta r \cos \theta.$$

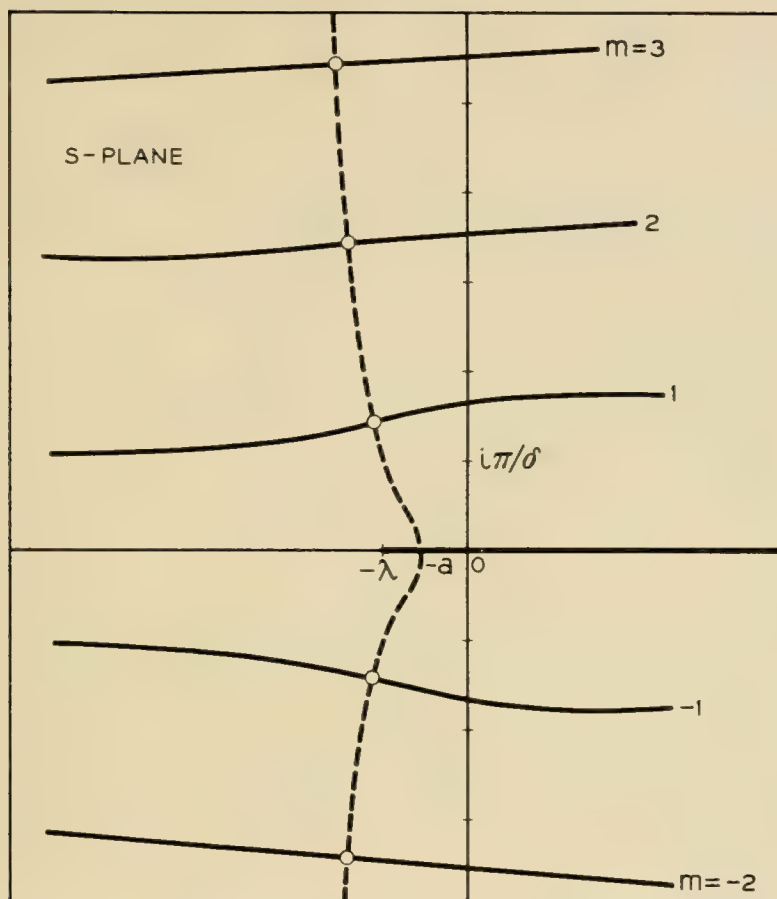
The first equation defines an infinite family of curves in the s -plane (see Fig. 5). The second equation defines a single curve which intersects the family at poles of $p(s)$.

2.5 Bounds on $F_0(L)$

As in Part I, we may derive bounds on $F_0(L)$ from the integral equation, and obtain

$$\left(1 - \frac{a}{\lambda}\right) e^{-aL} e^{-(a-\lambda)\delta} \leq F_0(L) \leq e^{-aL} e^{-(a-\lambda)\delta}.$$

Since $a = \lambda^2\delta + O(\delta^2)$ for small δ , the bounds are very close if $\lambda\delta$ is not too large.

FIG. 5—Solution of $s + \lambda = \lambda e^{-(s+\lambda)\delta}$

III COINCIDENCES BETWEEN n POISSON PATTERNS

3.1 Integral Equation

In this part we consider n one-dimensional Poisson patterns and ask for the probability, $F(L)$, that in the interval $(0, L)$ no pair of points from different patterns are coincident. Unlike Part I, we now consider only the case in which all n patterns have the same density λ . Let $P(L)$ be the conditional probability, given that Pattern No. 1 has a point at L , that there are no coincidences in $(0, L)$.

If $0 \leq L \leq \delta$, $P(L) = \exp - (n-1)\lambda L$.

If $\delta < L$,

$$P(L) = e^{-(n-1)\lambda L} \left(1 + (n-1)\lambda e^{-\lambda\delta} \int_0^{L-\delta} e^{(n-1)\lambda y} P(y) dy \right)$$

by the same sort of argument used in Part I. Then $F(L)$ will be given by

$$F(L) = e^{-n\lambda L} \left(1 + n\lambda \int_0^L e^{n\lambda x} P(x) dx \right).$$

3.2 Bounds and Asymptotic Formula

The Laplace transform of $P(L)$ is

$$p(s) = \{s + (n-1)\lambda(1 - e^{-(n\lambda+s)\delta})\}^{-1} \quad (3-1)$$

which has one real pole at a negative point $s = -a$, $a < (n-1)\lambda$. Again it is this pole which contributes the dominant term to both $P(L)$ and $F(L)$ for large L . We find

$$F(L) \approx \frac{n\lambda e^{-aL}}{(1 + [(n-1)\lambda - a]\delta)(n\lambda - a)}.$$

To bound $P(L)$ by expressions of the form $A \exp(-aL)$ one finds that $A > 1$ will give an upper bound and

$$A < \left(1 - \frac{a}{(n-1)\lambda}\right) e^{\lambda\delta}$$

will give a lower bound. The corresponding bounds on $F(L)$ are of the form

$$\left(1 - \frac{n\lambda A}{n\lambda - a}\right) e^{-n\lambda L} + \frac{n\lambda A}{n\lambda - a} e^{-aL}.$$

3.3 Exact Solution

As in Part II an exact formula for $F(L)$ may be given as a finite sum. We now derive it from the Laplace transform,

$$f(s) = (s + n\lambda)^{-1} (1 + n\lambda p(s)),$$

of $F(L)$. We may use (3-1) to expand $f(s)$ into the series

$$f(s) = \frac{1}{s + n\lambda} \left\{ 1 + n\lambda \sum_{k=0}^{\infty} \frac{((n-1)\lambda e^{-(n\lambda+s)\delta})^k}{(s + (n-1)\lambda)^{k+1}} \right\}. \quad (3-2)$$

The identity

$$\begin{aligned} & (s + n\lambda)^{-1} (s + (n-1)\lambda)^{-k-1} \\ &= \frac{1}{\lambda} \sum_{j=0}^k (-\lambda)^{-k+j} (s + (n-1)\lambda)^{-j-1} + (-\lambda)^{-k-1} (s + n\lambda)^{-1} \end{aligned} \quad (3-3)$$

provides a partial fraction expansion for the k^{th} term of the series (3-2). Transforming (3-2) term by term with the help of (3-3) we find

$$\begin{aligned} F(L) &= e^{-n\lambda L} [-(n-1)]^{[L/\delta]+1} \\ &+ ne^{-(n-1)\lambda L} \sum_{k=0}^{[L/\delta]} [-(n-1)e^{-\lambda\delta}]^k \sum_{j=0}^k \frac{[-\lambda(L-k\delta)]^j}{j!}. \end{aligned}$$

This is the desired formula for $F(L)$.

IV MULTIDIMENSIONAL PROBLEMS

4.1 Two-Pattern Lower Bound

We now derive some results on the probabilities of no coincidences in some multi-dimensional situations. The simplest one is a lower bound for the case of two Poisson patterns.

Theorem: Consider a d -dimensional region of volume V containing two Poisson patterns with densities λ and μ . Let $S(\delta)$ be the volume of the d -dimensional sphere of radius δ . The probability of no coincidences between the two patterns has the lower bound

$$e^{-\lambda V (1 - e^{-\mu S(\delta)})}.$$

Proof

Let the pattern with density λ be called the λ -pattern and the other the μ -pattern. Given any λ -pattern of k points there will be no coincidences provided only that a certain region T contains no points of the μ -pattern. T consists of all points of the volume V which lie in any of the spheres of radius δ centered on the k points of the λ -pattern. Since these spheres may overlap and may extend partly outside the volume V , we have

$$\text{volume of } T \leq k S(\delta),$$

and

$$\begin{aligned} \text{Prob (no coinc., given } k \text{ points)} &= \exp (-\mu \text{ volume of } T) \\ &\geq \exp (-k\mu S(\delta)). \end{aligned}$$

Since the number, k , of points of the λ -pattern has the Poisson distribution with mean λV the (unconditional) probability of no coincidences has the lower bound

$$\sum_{k=0}^{\infty} \frac{(\lambda V)^k}{k!} e^{-\lambda V} e^{-k\mu S(\delta)}.$$

Summing the series one proves the theorem. Interchanging λ and μ in the theorem gives another lower bound. The one stated above is the better of the two if $\lambda < \mu$.

The difference between the lower bound and the true probability comes from two sources: (a) The overlap between the k spheres; this will be a small effect if $\lambda^2 S(2\delta)V$ is small, and (b) the spheres which extend partly outside the volume V ; there will be relatively few such spheres if only a small fraction of the volume V lies within distance δ

of its boundary. Hence in some cases the lower bound will be a good approximation to the correct value.

It may also be noted that no real use was made of the spherical shape of the volumes $S(\delta)$. If one wants to consider a point of the μ -pattern to be coincident with a point of the λ -pattern if it lies in some other neighborhood, not of spherical shape, the same lower bound applies but with $S(\delta)$ replaced by the volume of the neighborhood.

4.2 Single-Pattern Lower Bound

A similar derivation in the case of a single Poisson pattern leads to:

Theorem: Let a Poisson pattern of density λ be distributed over a d -dimensional region of volume V . Let $S(\delta)$ be the volume of the d -dimensional sphere of radius δ . Then the probability of no coincidences is at least as large as

$$e^{-\lambda V} \{1 + \lambda S(\delta)\}^{V/S(\delta)}.$$

The theorem will follow from another bound which is slightly more accurate but much more cumbersome.

Lemma

In the above theorem a lower bound is

$$e^{-\lambda V} \left(1 + \lambda V + \sum_{k=2}^{\lfloor V/S(\delta) \rfloor} \frac{(\lambda V)^k}{k!} \prod_{j=1}^{k-1} [1 - jS(\delta)/V] \right). \quad (4-1)$$

Proof of Lemma

The probability sought is of the form

$$\sum_k e^{-\lambda V} \frac{(\lambda V)^k}{k!} p_k, \quad (4-2)$$

where p_k is the probability that, when exactly k points are distributed at random over V , there are no coincidences. To estimate p_k , imagine the k points to be numbered $1, 2, \dots, k$ and placed in the region one at a time. If no coincidences have been created among points $1, \dots, j$ (which is an event of probability p_j) the probability that the addition of point $j + 1$ creates no coincidence is just the probability that this new point lies in none of the j spheres of radius δ centered on points $1, \dots, j$. The union of these j spheres intersected with the volume V

is always of volume $\leq jS(\delta)$. Hence

$$p_{j+1} \geq p_j[1 - jS(\delta)/V],$$

or

$$p_k \geq \prod_{j=1}^{k-1} [1 - jS(\delta)/V]. \quad (4-3)$$

When $(k - 1)S(\delta) > V$ the above argument fails because the later terms of the product are negative; in this case we use the trivial bound $p_k \geq 0$. Combining (4-2) with (4-3) the lemma follows.

Once more the bound may be expected to be almost correct if $\lambda^2 VS(2\delta)$ is small and if most of the region V lies farther than δ away from its boundary. The bound is also correct for non-spherical neighborhoods (see discussion of previous theorem).

When $V/S(\delta)$ is large, the sum (4-1) is unwieldy. If we let H equal $V/S(\delta)$, we may rewrite the typical term in the sum as

$$\frac{(\lambda V)^k}{k!} \prod_{j=1}^{k-1} (1 - j/H) = \frac{(\lambda V/H)^k}{k!} H(H - 1) \cdots (H - k + 1).$$

If H happens to be an integer, this equals

$$\binom{H}{k} (\lambda V/H)^k,$$

so that the complete sum (4-1) equals

$$e^{-\lambda V} \left(1 - \frac{\lambda V}{H}\right)^H. \quad (4-4)$$

We will now prove that if H is not an integer, the sum always *exceeds* (4-4), so that (4-4) is a lower bound in all cases. We wish to prove that

$$1 + \sum_{k=1}^{[H]+1} \frac{x^k}{k!} H(H - 1) \cdots (H - k + 1) \geq (1 + x)^H \quad (4-5)$$

for any positive H , in which event the theorem follows with

$$x = \frac{\lambda V}{H} \quad \text{and} \quad H = V/S(\delta).$$

The inequality (4-5) will be proved by induction on $[H]$. If $[H] = 0$, then we are required to show that

$$1 + Hx \geq (1 + x)^H$$

for $0 \leq H < 1$. This follows immediately from the concavity of $(1+x)^H$.

Suppose now that (4-5) holds for a value H . If we integrate both sides of (4-5) from 0 to x , we obtain

$$x + \sum_{k=1}^{[H]+1} \frac{x^{k+1}}{(k+1)!} H(H-1) \cdots (H-k+1) \geq \frac{(1+x)^{H+1} - 1}{H+1},$$

which may be rewritten as

$$1 + \sum_{k=1}^{[H+1]+1} \frac{x^k}{k!} (H+1)(H) \cdots (H-k+2) \geq (1+x)^{H+1}.$$

This completes the induction, and the proof of the theorem.

4.3 Another Lower Bound (Any Number of Patterns)

Another kind of lower bound can be derived which sometimes will be better than the above bounds when the region V has a large fraction of its volume within δ of the boundary. For example, V might be a three-dimensional circular cylinder (a cable) with a radius which is comparable to δ .

To derive this bound one first finds the expected number, E , of coincidences in V . An upper bound on E will also suffice. Then it is noted that $1 - E$ is a lower bound on the probability of no coincidences. For if Q_N is the probability of finding N coincidences,

$$E = \sum N Q_N \geq \sum_{N=1}^{\infty} Q_N = 1 - Q_0. \quad (4-6)$$

4.4 Thick Cable

For example, we now give a lower bound which is of interest in connection with the problem of a cable with many wires.

Theorem: Let a Poisson pattern of points with density λ be placed in a cylinder of length L and radius $R > \delta$. The probability of finding no coincidences in the cylinder is at least as great as

$$1 - \lambda^2 \pi^2 L \left(\frac{2R^2 \delta^3}{3} - \frac{R \delta^4}{4} + \frac{\delta^5}{15} \right).$$

Proof

Introduce cylindrical coordinates r, φ, Z so that the cylinder is described by

$$r \leq R, \quad 0 \leq Z \leq L.$$

Consider first any pattern point (r, φ, Z) with Z -coordinate satisfying $\delta \leq Z \leq L - \delta$. Let arrows be drawn from this point to all other pattern points (if any) within distance δ . The expected number of arrows drawn from this point will be $\lambda G(r)$ where $G(r)$ is the volume of the intersection of the cylinder with a sphere of radius δ centered at the point. For points near the ends of the cylinder ($Z \leq \delta$ or $L - \delta \leq Z$), the expected number of arrows will be less than $\lambda G(r)$. Since the probability of finding a pattern point in a little volume element dV is λdV , we conclude that the expected number of arrows drawn in the entire cylinder will be less than

$$\iiint_{\text{cylinder}} \lambda^2 G(r) dV.$$

If the cylinder has N coincidences, there will be $2N$ arrows (each point of a coincident pair appears once at the head of an arrow and once at the tail). Hence the expected number of coincidences is

$$E \leq \lambda^2 \pi L \int_0^R G(r) r dr. \tag{4-7}$$

Since an exact formula for $G(r)$ is rather cumbersome, we are content with a simple but close upper bound. If $r \leq R - \delta$ then clearly $G(r) = 4\pi\delta^3/3$. If $r > R - \delta$ we get an upper bound on $G(r)$ by computing the shaded volume in Fig. 6; the intersection of the sphere with a half-space.

$$G(r) \leq [2\delta^3 + 3(R - r)\delta^2 - (R - r)^3] \pi/3.$$

Substituting these expressions for $G(r)$ in (4-7), integrating, and using (4-6) the theorem follows.

The approximation to $G(r)$ which was made above is bad when R is much less than δ , but in this case good estimates may be obtained from the one-dimensional results of Part II. Note also that if λ is large enough, the bound becomes negative and is therefore useless.

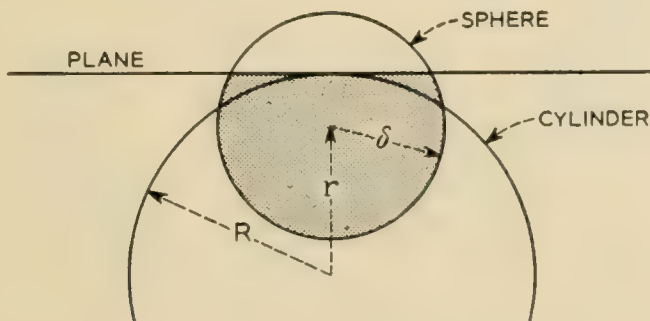


Fig. 6 — A region for estimating $G(r)$.

4.5 Upper Bounds

Good upper bounds appear even harder to get than lower bounds. One procedure is to divide the region V into a number of smaller cells. If each cell has probability, p , of no coincidences and if there are K cells, then p^K is the probability of no coincidence in any cell. If there is no coincidence in V there will be none in any cell; hence p^K is an upper bound on the probability of no coincidence in V .

Of course, p^K is too large because of the possibility of a coincidence between two points in different cells. It follows that p^K will be a close bound only if the cell size is made large; but then p becomes hard to compute.

For example, consider self-coincidences in a single Poisson pattern in a large region of area V in the plane. Cover this area with an array of hexagonal cells of side $\delta/2$ as shown in Fig. 7. The area of each hexagon is $3\sqrt{3}\delta^2/8$ so the number of cells used will be about $K = 8V/3\sqrt{3}\delta^2$. A cell has no coincidence if it contains at most one pattern point, hence

$$p = (1 + \lambda 3\sqrt{3}\delta^2/8) \exp - 3\sqrt{3}\lambda\delta^2/8.$$

The upper bound is

$$p^K = e^{-\lambda V} \left(1 + \frac{3\sqrt{3}}{8} \lambda \delta^2 \right)^{(8V/3\sqrt{3}\delta^2)}$$

which has an interesting resemblance to the lower bound

$$e^{-\lambda V} (1 + \pi \lambda \delta^2)^{V/\pi \delta^2}.$$

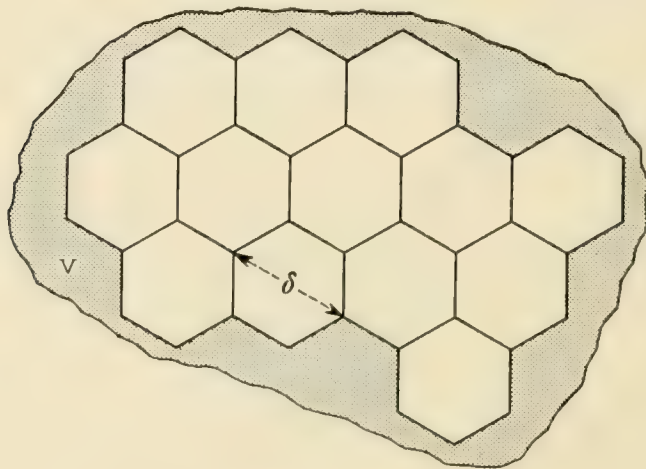


Fig. 7 — Pattern for studying coincidences in a plane region.

4.6 An Exact Calculation

The upper and lower bounds in Section 4.5 are not very close, largely because of the small size of the hexagonal cells. An improved upper bound may be obtained using square cells of side 2δ . We can calculate p for small rectangular cells but only if we redefine our notion of coincidence in terms of square neighborhoods instead of circular neighborhoods. That is, points (x_1, y_1) and (x_2, y_2) are now considered coincident if simultaneously

$$|x_1 - x_2| \leq \delta, \quad \text{and} \quad |y_1 - y_2| \leq \delta.$$

The result we get is the only exact calculation of a non-trivial multi-dimensional coincidence probability known to us.

Consider the rectangle $0 \leq x \leq L$, $0 \leq y \leq M$ with L and M both $\leq 2\delta$. If L is less than δ , two points are coincident if and only if their y -coordinates differ by less than δ . The problem then reduces to a one-dimensional coincidence computation such as we gave in Part II. Therefore, suppose both L and M are greater than δ .

There is probability

$$g_k = \frac{(\lambda LM)^k}{k!} e^{-\lambda LM}$$

that the rectangle contains k points. We therefore subdivide the problem into cases of the form "given k , find the probability that the k points have no coincidences". Only five of these cases have a non-zero answer. To show this, divide the rectangle into four rectangles of sides $L/2$, $M/2$; if $k \geq 5$ one of these rectangles must contain more than one point, and so a coincidence. The remaining cases $k = 0, 1, 2, 3, 4$ may be further subdivided according to which pairs of x -coordinates are less than δ apart. Let us number the k points $(x_1, y_1), \dots, (x_k, y_k)$ in such a way that the x -coordinates are in order $x_1 \leq x_2 \leq \dots \leq x_k$. If, for some i , $x_{i+2} \leq x_i + \delta$, then the subcase in question contributes zero to the probability of no coincidences because all of $|x_i - x_{i+1}|$, $|x_{i+1} - x_{i+2}|$, $|x_{i+2} - x_i|$ are $\leq \delta$ and at least one of $|y_i - y_{i+1}|$, $|y_{i+1} - y_{i+2}|$, $|y_{i+2} - y_i|$ is $\leq \delta$. The only subcases which remain to give a non-zero contribution are the nine listed in Table I. The number in the "subcase" column is k . The next column contains the x -inequalities which define the subcase. The probability that the k ordered x -coordinates satisfy the stated inequalities is listed as prob_x . If the x -inequalities are satisfied there will be no coincidences if and only if $|y_b - y_a| > \delta$ for every inequality $|x_b - x_a| \leq \delta$ given in the x -inequality column. These y -in-

TABLE I

Subcase	x inequ.	prob_x	y inequ.	prob_y
0	—	1	—	1
1	—	1	—	1
2(a)	$x_2 - x_1 > \delta$	$(1 - \delta/L)^2$	—	1
2(b)	$x_2 - x_1 \leq \delta$	$\frac{2\delta L - \delta^2}{L^2}$	$ y_2 - y_1 > \delta$	$(1 - \delta/M)^2$
3(a)	$x_2 - x_1 \leq \delta$ $x_3 - x_2 > \delta$	$\frac{2}{3}(1 - \delta/L)^3$	$ y_2 - y_1 > \delta$	$(1 - \delta/M)^2$
3(b)	$x_2 - x_1 > \delta$	$\frac{2}{3}(1 - \delta/L)^3$	$ y_2 - y_3 > \delta$	$(1 - \delta/M)^2$
3(c)	$x_2 - x_1 \leq \delta$ $x_3 - x_2 \leq \delta$ $x_3 - x_1 > \delta$	$\frac{2}{3}(1 - \delta/L)^2 \left(\frac{4\delta}{L} - 1 \right)$	$ y_2 - y_3 > \delta$ $ y_1 - y_2 > \delta$	$\frac{2}{3}(1 - \delta/M)^3$
4(a)	$x_3 - x_2 \leq \delta$ $x_3 - x_1 > \delta$ $x_4 - x_2 > \delta$	$\frac{1}{3}(1 - \delta/L)^4$	$ y_2 - y_1 > \delta$ $ y_3 - y_2 > \delta$ $ y_4 - y_3 > \delta$	$\frac{5}{12}(1 - \delta/M)^4$
4(b)	$x_3 - x_2 > \delta$	$\frac{1}{3}(1 - \delta/L)^4$	$ y_2 - y_1 > \delta$ $ y_4 - y_3 > \delta$	$(1 - \delta/M)^4$

equalities are listed in the third column and the probabilities that they are satisfied are listed as prob_y . The probability of no coincidences is

$$\sum g_k \text{prob}_x \text{prob}_y$$

where the sum is over all nine subcases. The sum is

$$\begin{aligned} \exp(-\lambda LM) \left\{ 1 + \lambda LM + \frac{\lambda^2}{2} [L^2 M^2 - \delta^2 (2L - \delta)(2M - \delta)] \right. \\ + \frac{2\lambda^3}{27} (L - \delta)^2 (M - \delta)^2 (2LM + L\delta + M\delta - 4\delta^2) \\ \left. + \frac{17\lambda^4}{864} (L - \delta)^4 (M - \delta)^4 \right\}. \end{aligned}$$

If $L = M = 2\delta$, this reduces to

$$\exp(-4\delta^2\lambda) \left[1 + 4\delta^2\lambda + \frac{7}{2}\delta^4\lambda^2 + \frac{16}{27}\delta^6\lambda^3 + \frac{17}{864}\delta^8\lambda^4 \right].$$

A sample of one of the above computations may be instructive. Consider, for example, Case 4(a). We have $0 \leq x_1 \leq x_2 \leq x_3 \leq x_4 \leq L$, and require:

$$x_3 - x_2 \leq \delta,$$

$$x_3 - x_1 > \delta,$$

$$x_4 - x_2 > \delta.$$

The probability of this is

$$\begin{aligned} (L^4/8)^{-1} \int_{x_3=\delta}^L \int_{x_2=x_3-\delta}^{L-\delta} \int_{x_1=0}^{x_3-\delta} \int_{x_4=x_2+\delta}^L dx_4 dx_1 dx_2 dx_3 \\ = \frac{8}{L^4} \int_{x_3=\delta}^L \int_{x_2=x_3-\delta}^{L-\delta} (L - x_2 - \delta)(x_3 - \delta) dx_3 dx_2 \\ = \frac{8}{L^4} \frac{(L - \delta)^4}{24} = \frac{1}{3} \left(1 - \frac{\delta}{L}\right)^4. \end{aligned}$$

In the y -direction we require $|y_2 - y_1| > \delta$, $|y_3 - y_2| > \delta$, $|y_4 - y_3| > \delta$, and there are no order restrictions. Assume first that $y_2 < y_3$. Then the probability that y_1 and y_4 satisfy their restrictions is

$$\left(\frac{y_3 - \delta}{M}\right) \left(\frac{M - y_2 - \delta}{M}\right).$$

Hence, the probability for satisfying all the conditions is

$$\int_{\delta}^M \int_0^{y_3-\delta} \left(\frac{y_3 - \delta}{M}\right) \left(\frac{M - y_2 - \delta}{M}\right) \frac{dy_2}{M} \frac{dy_3}{M} = \frac{5}{24} \left(1 - \frac{\delta}{M}\right)^4.$$

Interchanging y_2 with y_3 and y_1 with y_4 shows that the assumption $y_2 > y_3$ yields the same answer, so that the required probability is

$$\frac{5}{12} \left(1 - \frac{\delta}{M}\right)^4.$$

V NUMERICAL WORK

5.1 Coincidences between Two Patterns

5.1.1 Machine Computation of $F(L)$

To compute the probability of no coincidences in a line of length L directly, it is convenient to transform equations (1-2) through (1-4) into

the following differential difference equations:

$$P_1'(x) + \mu P_1(x) = \begin{cases} 0 & \text{if } x \leq \delta \\ P_2(x - \delta)\mu e^{-(\lambda+\mu)\delta} & \text{if } x > \delta, \end{cases}$$

$$P_2'(x) + \lambda P_2(x) = \begin{cases} 0 & \text{if } x \leq \delta \\ P_1(x - \delta)\lambda e^{-(\lambda+\mu)\delta} & \text{if } x > \delta, \end{cases}$$

$$F'(x) + (\lambda + \mu)F(x) = \lambda P_1(x) + \mu P_2(x),$$

$$P_1(0) = P_2(0) = F(0) = 1.$$

These have been solved on a general purpose analog computer with the aid of a lumped-element approximate delay line for a number of cases. We have chosen for illustrative purposes the parameters $\lambda = 5$, $\mu = 10$, $\delta = 0.02$, and $L \leq 1$. The exact solution, together with various approximations to be described in the sequel, is plotted in Fig. 8, where the exact solution is labelled y_1 .

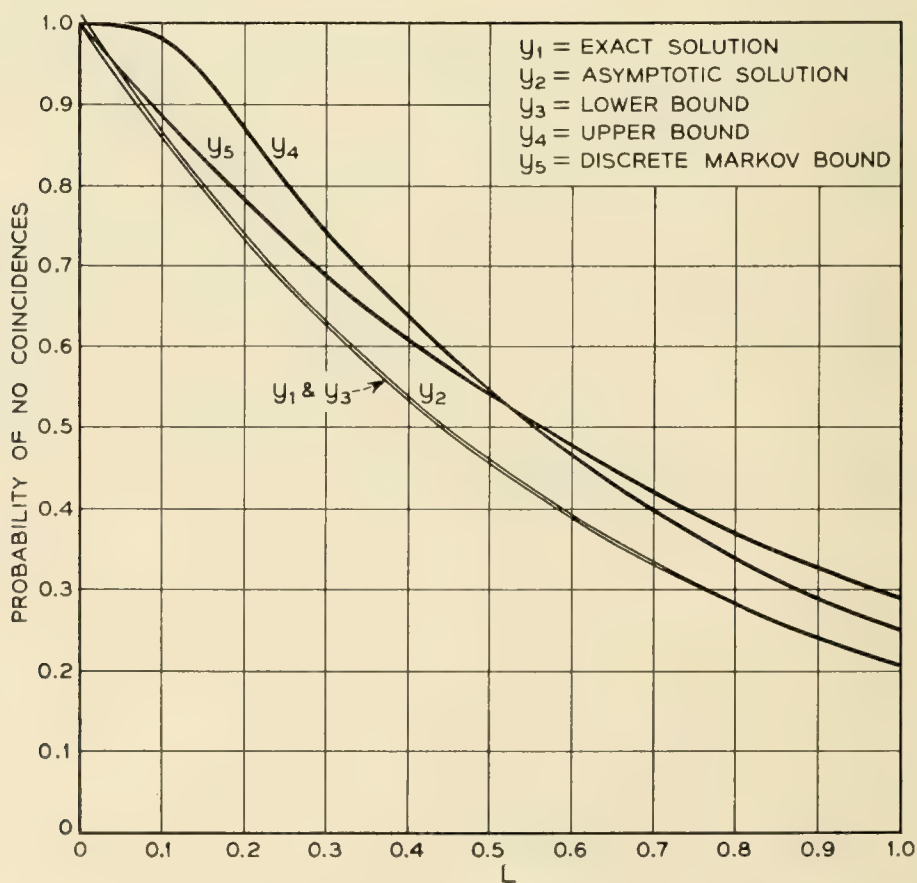


FIG. 8—Probability of no coincidences between two one-dimensional Poisson patterns with $\lambda = 5$, $\mu = 10$, if $\delta = 0.02$.

5.1.2 The Asymptotic Formula

An approximation to the probability $F(x)$ of no coincidences is given by the asymptotic formula (1-10) which, of course, becomes a better approximation the larger L becomes. If $\lambda = 5$, $\mu = 10$, and $\delta = 0.02$, the smallest value, a , such that

$$(\lambda - a)(\mu - a) = \lambda\mu e^{-2(\lambda+\mu-a)\delta}$$

is $a = 1.548$. The asymptotic formula for $F(L)$ now becomes

$$F(L) \approx 1.013e^{-1.548L},$$

which is found in Fig. 8 as y_2 .

5.1.3 Bounds Using the Asymptotic Exponent

Formulas (1-12) through (1-18) give a scheme for computing both upper and lower bounds for $F(L)$ which have the right behavior for large L , and also agree with the solution at $L = 0$. They become

$$F(L) \geq 1.007e^{-1.548L} - 0.007e^{-15L},$$

and

$$F(L) \leq 1.195e^{-1.548L} - 0.195e^{-15L},$$

respectively, and are represented by y_3 and y_4 in Figure 8.

5.1.4 An Upper Bound by a Discrete Markov Process

If we mark on the positive x -axis the points $n\delta/2$, $n = 0, 1, 2, \dots$, we can assign to each interval of length $\delta/2$ thus created a state (ij) , $i, j = 0$ or 1 , as follows: $i = 0$ if no point of the λ -process is present in the interval, $i = 1$ if one or more points of the λ -process are present, and similarly for j and μ . An interval of length δ , made up of two adjacent intervals of length $\delta/2$, may then be represented by a number between 0 and 15 in binary notation, where 3, 6, 7, 9, and 11–15 represent a coincidence within the interval of length δ . We now define a Markov process as follows: in the interval $0 \leq t < \delta$, let $p_i^{(0)}$, $i = 0, 1, 2, 4, 5, 8, 10$, be the probabilities of occurrence of the i^{th} state, so that, for example, $p_0^{(0)} = e^{-2\lambda\delta}e^{-2\mu\delta}$, and $p_1^{(0)} = e^{-2\lambda\delta}e^{-\mu\delta}(1 - e^{-\mu\delta})$. These are the states in which there is no coincidence in $(0, \delta)$. In addition, let $q^{(0)}$ represent the probability of all the other states put together; i.e., of a coincidence in $(0, \delta)$. We now define $p_i^{(n)}$, $i = 0, 1, 2, 4, 5, 8, 10$ as the probability of the i^{th} state in the interval $(n\delta/2, (n+2)\delta/2)$, where we

require in addition that all states in the intervals $(k\delta/2, (k+2)\delta/2)$, $k < n$, are from the same "no coincidence" index set. We define $q^{(n)}$ as the probability of a state 3, 6, 7, 9, or 11-15, in *some* interval $(k\delta/2, (k+2)\delta/2)$, $k \leq n$. There are then transition probabilities from states in the $n-1^{\text{st}}$ to states in the n^{th} interval. For example,

$$p_0^{(n)} = e^{-\lambda\delta} e^{-\mu\delta} (p_0^{(n-1)} + p_4^{(n-1)} + p_8^{(n-1)}),$$

and

$$q^{(n)} = q^{(n-1)} + (1 - e^{-\lambda\delta})(1 - e^{-\mu\delta})(p_0^{(n-1)} + p_4^{(n-1)} + p_8^{(n-1)}) \\ + (1 - e^{-\lambda\delta})(p_1^{(n-1)} + p_5^{(n-1)}) + (1 - e^{-\mu\delta})(p_2^{(n-1)} + p_{10}^{(n-1)}).$$

The quantity $1 - q^{(n)}$ is then an upper bound for the probability of no coincidences (upper because it is possible for a coincidence to occur in the process which is not counted in this subdivision of it). The curve y_5 in Fig. 8 is drawn through points at $L = n\delta/2$ computed in this manner.

To summarize the results, we see that the asymptotic formula and the lower bound are both indistinguishable from the right answer; the upper bounds are fairly far off. The upper bound derived by the Markov process is better than that derived from the integral equation until

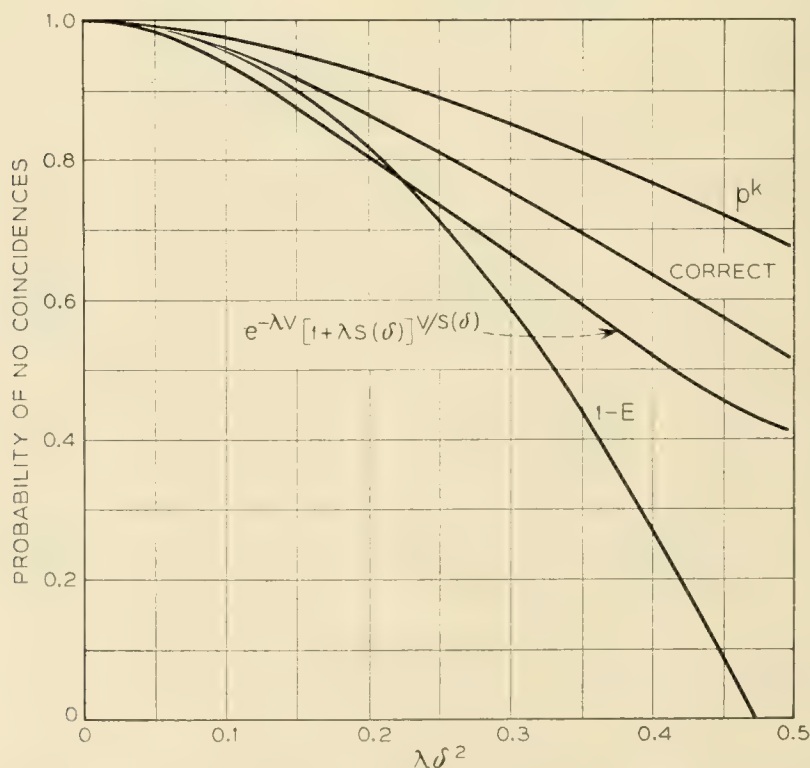


FIG. 9 — Probability of no coincidences in a $2\delta \times 2\delta$ square; neighborhoods are square.

about $L = 0.5$ (25 iterations), when the integral equation upper bound becomes better.

5.2 A Single Pattern in a Square

To test our higher-dimensional bounds, we consider again coincidences in a single Poisson pattern in a square of side 2δ . The exact probability of no coincidences was given in Part IV assuming square neighborhoods. The lower bound (Sec. 4.2)

$$e^{-\lambda V}(1 + \lambda S(\delta))^{V/S(\delta)}$$

applies using $V = (2\delta)^2$ and $S(\delta) = (2\delta)^2$ for square neighborhoods. To use the lower bound $1 - E$ we note that the exact expected number of coincidences is

$$E = \frac{1}{2} \lambda^2 \int_0^{2\delta} \int_0^{2\delta} A(x, y) \, dx \, dy$$

where $A(x, y)$ is the area of the intersection of the given square with the square neighborhood centered at (x, y) . The lower bound is $1 - E = 1 - 9\lambda^2\delta^4/2$. The upper bound p^K can be used if the square is cut into $K = 4$ squares of side δ , each with a probability $p = (1 + \lambda\delta^2) \exp - \lambda\delta^2$ of no coincidence.

These bounds, together with the exact probability, are plotted as functions of $\lambda\delta^2$ in Fig. 9. When $\lambda\delta^2$ is small, the $1 - E$ bound is correct to terms of order $O(\lambda^3\delta^6)$. This might have been predicted from (4-6) since it seems reasonable that $Q_2, Q_3, \dots$ should be of higher order in λ than Q_1 when λ is small. Ultimately the first lower bound becomes a better estimate. It must be recognized that this other lower bound is being tested under very severe conditions. Since every point of the square has a neighborhood which intersects the boundary, the errors from source (b) of Part V are considerable.

The authors wish to thank D. W. Hagelbarger and H. T. O'Neil for their assistance in the course of the calculations reported in this section, and Miss D. T. Angell for preparing some of the figures.

REFERENCES

1. C. Domb, The Problem of Random Intervals on a Line, Proc. Cambridge Phil. Soc., **43**, pp. 329-341, 1947.
2. P. Eggleton and W. O. Kermack, A Problem in the Random Distribution of Particles, Proc. Royal Soc. Edinburgh, Sec A, **62**, pp. 103-115, 1944.
3. W. Feller, On Probability Problems in the Theory of Counters, Studies and Essays presented to R. Courant, Interscience, pp. 105-115, 1948.
4. E. C. Molina, The Theory of Probability and Some Applications to Engineering Problems, Trans. A.I.E.E., **44**, pp. 294-299, 1925.
5. L. Silberstein, The Probable Number of Aggregates in Random Distributions of Points, Phil. Mag., Series 7, **36**, pp. 319-336, 1945.

Bell System Technical Papers Not Published in this Journal

ANDERSON, O. L., see Andreatch, P.

ANDREATCH, P., and ANDERSON, O. L.¹

Teflon as a Pressure Medium, Rev. Sci. Instr., **28**, p. 288, April, 1957.

BOYET, H.,¹ and SEIDEL, H.¹

Analysis of Nonreciprocal Effects in an N-Wire Ferrite-Loaded Transmission Line, Proc. I.R.E., **45**, pp. 491-495, April, 1957.

BOYET, H., see Seidel, H.

BOZORTH, R. M.¹

Review of Magnetic Annealing, Proc. of 1956 A.I.E.E. Conf. on Magnetism and Magnetic Materials, **T-91**, pp. 69-75, April, 1957.

BURNS, F. P.¹

Piezoresistive Semiconductor Microphone, J. Acous. Soc. Am., **29**, pp. 248-253, Feb., 1957.

CARRUTHERS, J. A., see Geballe, T. H.

CIOFFI, P. P.¹

Rectilinearity of Electron-Beam Focusing Fields from Transverse Component Determinations, Commun. and Electronics, **29**, pp. 15-19, March, 1957.

COOK, R. K.¹

Absorption of Sound by Patches of Absorbent Material, J. Acous. Soc. Am., **29**, pp. 324-329, March, 1957.

¹ Bell Telephone Laboratories, Inc.

COOK, R. K.¹

Variation of Elastic Constants and Static Strains with Hydrostatic Pressure: A Method for Calculation from Ultrasonic Measurements, J. Acous. Soc. Am., **29**, pp. 445-449, April, 1957.

DEWALD, J. F.¹

The Kinetics of Formation of Anode Films on Single Crystal Indium Antimonide, J. Electrochem. Soc., **104**, pp. 244-251, April, 1957.

DOWLING, R. C.⁴

Lightning Protection on the Stevens Point-Wisconsin Rapids Inter-city Telephone Cable, Commun. and Electronics, **28**, pp. 697-701, Jan., 1957.

FEHER, G.¹

Electron Structure of F Centers in KCl by the Electron Spin Double-Resistance Technique, Phys. Rev., Letter to the Editor, **105**, pp. 1122-1123, Feb. 1, 1957.

FOSTER, F. G.¹

The Unconventional Application of the Metallograph, Focus, **18**, pp. 16-20, April, 1957.

GARN, P. D.¹

An Automatic Recording Balance, Anal. Chem., **29**, pp. 839-841, May, 1957.

GAST, R. W.⁵

Field Experience with the A2A Video System, Commun. and Electronics, **28**, pp. 710-716, Jan., 1957.

GEBALLE, T. H.,¹ CARRUTHERS, J. A.,⁶ ROSENBERG, H. M.,⁶ and ZIMAN, J. M.⁷

The Thermal Conductivity of Germanium and Silicon Between 2 and 300°K, Proc. Royal Soc., **A238**, pp. 502-514, Jan. 29, 1957.

¹ Bell Telephone Laboratories, Inc.

⁴ Wisconsin Telephone Company, Madison.

⁵ New York Telephone Company, New York.

⁶ Oxford University, England.

⁷ Cambridge University, England.

GELLER, S.¹

Crystallographic Studies of Perovskite-Like Compounds.—IV. Rare Earth Scandates, Vanadites, Galliates, Orthochromites, Acta Crys., 10, pp. 243–248, April 10, 1957.

GELLER, S.¹

Crystallographic Studies of Perovskite-Like Compounds.—V. Relative Ionic Sizes, Acta Crys., 10, pp. 248–251, April 10, 1957.

GELLER, S.,¹ and GILLES, M. A.¹

Structure and Ferrimagnetism of Ythium and Rare-Earth-Iron Garnets, Acta Crys., 10, p. 239, March 10, 1957.

GILLES, M. A.¹

Crystallographic Studies of Perovskite-Like Compounds.—III. $\text{La}(\text{M}_x, \text{Mn}_{1-x})\text{O}$ with $\text{M} = \text{Co}, \text{Fe}$ and Cr , Acta Crys., 10, pp. 161–166, March, 1957.

GILLES, M. A., see Geller, S.

GREEN, E. I.¹

Nature's Pulses, I.R.E. Student Quarterly, 3, pp. 3–5, Feb., 1957.

GROSSMAN, A. J.¹

Synthesis of Tschebycheff Parameter Symmetrical Filters, Proc. I.R.E., 45, pp. 454–473, April, 1957.

HAGSTRUM, H. D.¹

Thermionic Constants and Sorption Properties of Hafnium, J. Appl. Phys., 28, pp. 323–328, March, 1957.

HAWORTH, F. E.¹

Breakdown Fields of Activated Electrical Contacts, J. Appl. Phys., Letter to the Editor, 28, p. 381, March, 1957.

HEIDENREICH, R. D.,¹ and NESBITT, E. A.¹

Stacking Disorders in Nickel Base Magnetic Alloys, Phys. Rev., Letter to the Editor, 105, pp. 1678–1679, March 1, 1957.

¹ Bell Telephone Laboratories, Inc

HOLDEN, A. N., see Wood, Elizabeth A.

HUNTLEY, H. R.²

Where We Are and Where We Are Going in Telephone Transmission, *Commun. and Electronics*, **29**, pp. 54-63, March, 1957.

JOEL, A. E.¹

Electronics in Telephone Switching Systems, *Commun. and Electronic*, **28**, pp. 701-709, Jan., 1957.

KAPPEL, F. R.²

Three-Dimensional Engineers, *Elec. Engg.*, **76**, pp. 267-270, April 1957.

KARLIN, J. E., see Pierce, J. R.

KARP, A.¹

Backward-Wave Oscillator Experiments at 100 to 200 Kilomegacycles, *Proc. I.R.E.*, **45**, pp. 496-503, April, 1957.

KELLY, M. J.¹

The Work and Environment of the Physicist Yesterday, Today, and Tomorrow, *Phys. Today*, **10**, pp. 26-31, April, 1957.

LAW, J. T.,¹ and MEIGS, P. S.¹

Rates of Oxidation of Germanium, *J. Electrochem. Soc.*, **104**, pp. 154-159, March, 1957.

LAW, J. T.,¹ and MEIGS, P. S.¹

The High Temperature Oxidation of Germanium, Semiconductor Surface Physics (book), pp. 383-400, 1957, Univ. of Penna. Press, Philadelphia.

LIEHR, A. D.¹

Structure of $\text{Co}(\text{CO})_4\text{H}$ and $\text{Fe}(\text{CO})_4\text{H}_2$, *Zeitschrift Für Naturforschung*, **12b**, pp. 95-96, Feb., 1957.

¹ Bell Telephone Laboratories, Inc.

² American Telephone and Telegraph Company.

LIEHR, A. D.¹

Structure of π -Cyclopentadienyl Metal Hydrides, *Naturwissenschaften*, **44**, p. 61, Feb. 1, 1957.

LUNDBERG, C. V., see Vacca, G. N.

LUNDBERG, J. L.,¹ and NELSON, L. S.¹

The High Intensity Flash Irradiation of Polymers, *Nature*, Letter to the Editor, **179**, pp. 367-368, Feb. 16, 1957.

MARRISON, W. A.¹

A Wind-Operated Electric Power Supply, *Elec. Engg.*, **76**, pp. 418-421, May, 1957.

McCALL, D. W.¹

Nuclear Magnetic Resonance in Guanidinium Aluminum Sulfate Hexahydrate, *J. Chem. Phys.*, Letter to the Editor, **26**, pp. 706-707 March, 1957.

MEIGS, P. S., see Law, J. T.

MENDIZZA, A.,¹ SAMPLE, C. H.,⁸ and TEEL, R. B.⁹

A Comparison of the Corrosion Behavior and Protective Value of Electrodeposited Zinc and Cadmium on Steel, *Symposium on Properties, Tests, and Performances of Electrodeposited Metallic Coatings*, **A.S.T.M. Special Tech. Publication 197**, pp. 49-64, 1957.

MILLER, S. L.¹

The Ionization Rates for Holes and Electrons in Si, *Phys. Rev.*, **105**, pp. 1246-1249, Feb. 15, 1957.

NELSON, L. S., see Lundberg, J. L.

NESBITT, E. A., see Heidenreich, R. D.

¹ Bell Telephone Laboratories, Inc.

⁸ International Nickel Company, New York City.

⁹ International Nickel Company, Wrightsville Beach, North Carolina.

PALMQUIST, T. F.¹⁰

Multi-Unit Neutralizing Transformers, Commun. and Electronics, **28**, pp. 717-721, Jan., 1957.

PIERCE, J. R.,¹ and KARLIN, J. E.¹

Information Rate of a Human Channel, Proc. I.R.E., Letter to the Editor, **45**, p. 368, March, 1957.

REA, W. T.¹

The Communication Engineer's Needs in Information Theory, Commun. and Electronics, **28**, pp. 805-808, Jan., 1957.

ROSENBERG, H. M., see GEBALLE, T. H.

SAMPLE, C. H., see MENDIZZA, A.

SEIDEL, H.,¹ and BOYET, H.¹

Form of Polder Tensor for Single Crystal Ferrite with Small Cubic Symmetry Anisotropy Energy, J. Appl. Phys., **28**, pp. 452-454, April, 1957.

SEIDEL, H., see Boyet, H.

SLICHTER, W. P.¹

Nuclear Magnetic Resonance in Some Fluorine Derivatives of Polyethylene, J. Poly. Sci., **24**, pp. 173-188, April, 1957.

SMITH, K. D., see Veloric, H. S.

SUHL, H.¹

Proposal for a Ferromagnetic Amplifier in the Microwave Range, Phys. Rev., Letter to the Editor, **106**, pp. 384-385, April 15, 1957.

SWANEKAMP, F. W., see Van Uitert, L. G.

TEEL, R. B., see Mendizza, A.

¹ Bell Telephone Laboratories, Inc.

¹⁰ Bell Telephone Company of Canada, Montreal, Quebec.

TREUTING, R. G.¹

Torsional Strain and the Screw Dislocation in Whisker Crystals, *Acta Met.*, Letter to the Editor, **5**, pp. 173-175, March, 1957.

VACCA, G. N.,¹ and LUNDBERG, C. V.¹

Aging of Neoprene in a Weatherometer, *Wire and Wire Products*, **32**, pp. 418-457, April, 1957.

VAN UITERT, L. G.¹

Magnesium-Copper — Manganese-Aluminum Ferrites for Microwave Applications, *J. Appl. Phys.*, **28**, pp. 320-322, March, 1957.

VAN UITERT, L. G.¹

Magnetic Induction and Coercive Force Data on Members of the Series $\text{BaAl}_x\text{Fe}_{12-x}\text{O}_{19}$ and Related Oxides, *J. Appl. Phys.*, **28**, pp. 317-319, March, 1957.

VAN UITERT, L. G.,¹ and SWANEKAMP, F. W.¹

Permanent Magnet Oxides Containing Divalent Metal Ions, *J. Appl. Phys.*, **28**, pp. 482-485, April, 1957.

VELORIC, H. S.,¹ and SMITH, K. D.¹

Silicon Diffused Junction Avalanche Diodes, *J. Electrochem. Soc.*, **104**, pp. 222-227, April, 1957.

WALKER, L. R.¹

Orthogonality Relays for Gyrotropic Wave Guides, *J. Appl. Phys.*, Letter to the Editor, **28**, p. 377, March, 1957.

WEBER, L. A.¹

Influence of Noise on Telephone Signaling Circuit Performance, *Commun. and Electronics*, **28**, pp. 636-643, Jan., 1957.

WEIBEL, E. S.¹

An Electronic Analogue Multiplier, *Trans. I.R.E., PGEC, EC-6*, pp. 30-34, March, 1957.

¹ Bell Telephone Laboratories, Inc.

WEISS, J. A.¹

An Interference Effect Associated with Faraday Rotation, and Its Application to Microwave Switching, Proc. Conf. on Magnetism and Magnetic Materials, pp. 580-585, Feb., 1957.

WERTHEIM, G. K.¹

Energy Levels in Electron-Bombarded Silicon, Phys. Rev., **105**, pp. 1730-1735, March 15, 1957.

WILLIS, F. H.¹

Some Results with Frequency Diversity in a Microwave Radio System, Commun. and Electronics, **29**, pp. 63-67, March, 1957.

WOOD, ELIZABETH A.,¹ and HOLDEN, A. N.¹

Monoclinic Glycine Sulfate: Crystallographic Data, Acta Crys., **10**, pp. 145-146, Feb., 1957.

YOUNKER, E. L.¹

A Transistor Driven Magnetic Core Memory, Trans. I.R.E., PGEC, **EC-6**, pp. 14-20, March, 1957.

ZIMAN, J. M., see Geballe, T. H.

¹ Bell Telephone Laboratories, Inc.

Recent Monographs of Bell System Technical Papers Not Published in This Journal*

BASHKOW, T. R., and DESOER, C. A.

A Network Proof of a Theorem on Hurwitz Polynomials, Monograph 2614.

BEACH, A. L., see Thurmond, C. D.

BIGGS, B. S., see Lundberg, C. V.

CUTLER, C. C.

Instability in Hollow and Strip Electron Beams, Monograph 2711.

DESOER, C. A., see Bashkow, T. R.

ELIAS, P., FEINSTEIN, A., and SHANNON, C. E.

A Note on the Maximum Flow Through a Network, Monograph 2768.

FEINSTEIN, A., see Elias, P.

GULDNER, W. G., see Thurmond, C. D.

HARING, H. E., see Taylor, R. L.

INGRAM, S. B.

Role of Evening Engineering Education in the Training of Technicians, Monograph 2771.

KRAMER, H. P., and MATHEWS, M. V.

A Linear Coding for Transmitting a Set of Correlated Signals, Monograph 2757.

* Copies of these monographs may be obtained on request to the Publication Department, Bell Telephone Laboratories, Inc., 463 West Street, New York 14, N. Y. The numbers of the monographs should be given in all requests.

LEWIS, H. W., SMITH, DE WITT H., and LEWIS, M. R.

Ballistocardiographic Instrumentation, Monograph 2747.

LEWIS, M. R., see Lewis, H. W.

LUNDBERG, C. V., VACCA, G. N., and BIGGS, B. S.

Resistance of Rubber Compounds to Outdoor and Accelerated Ozone Attack, Monograph 2773.

MATHEWS, M. V., see Kramer, H. P.

McMILLAN, B.

Two Inequalities Implied by Unique Decipherability, Monograph 2774.

ROBERTSON, S. D.

Recent Advances in Finline Circuits, Monograph 2759.

SHANNON, C. E.

Zero Error Capacity of a Noisy Channel, Monograph 2760.

SHANNON, C. E., see Elias, P.

SMITH, DE WITT H., see Lewis, H. W.

TAYLOR, R. L., and HARING, H. E.

A Metal-Semiconductor Capacitor, Monograph 2776.

THEUERER, H. C.

Removal of Boron from Silicon by Hydrogen Water Vapor Treatment, Monograph 2762.

THURMOND, C. D., GULDNER, W. G., and BEACH, A. L.

Hydrogen and Oxygen in Single-Crystal Germanium Determined by Gas Analysis, Monograph 2777.

TRUMBORE, F. A.

Solid Solubilities and Electrical Properties of Tin in Germanium Single Crystals, Monograph 2779.

VACCA, G. N., see Lundberg, C. V.

Contributors to This Issue

VACLAV E. BENES, A.B., Harvard College, 1950; M.A. Ph.D., Princeton University, 1953. Instructor in logic and philosophy of science, Princeton University, 1952-53; Bell Telephone Laboratories, 1953-. Since joining the Laboratories, Mr. Benes has been engaged in mathematical systems research, involving stochastic processes describing the passage of traffic through a switching system. He is the author of a number of papers on mathematical logic and analytic philosophy. Member of the Mind Association, the Association for Symbolic Logic, the Institute of Mathematical Statistics, American Mathematical Society, and Phi Beta Kappa.

E. N. GILBERT, B.S., Queens College, 1942; Ph.D., Massachusetts Institute of Technology, 1948; Mr. Gilbert's early employment was with the M.I.T. Radiation Laboratory. He joined Bell Telephone Laboratories in 1948. His work has been on studies of the information theory and on the switching theory. He now is part of the communication fundamentals group. Mr. Gilbert is a member of the American Mathematical Society.

H. O. POLLAK, B.A., Yale University, 1947; M.A., Harvard University, 1948; Ph.D., Harvard University, 1951; Bell Telephone Laboratories, 1951-. Mr. Pollak has been engaged in mathematical research and military systems analysis. He is a member of Phi Beta Kappa, Sigma Xi, American Mathematical Society and Mathematical Association of America.

M. B. PRINCE, A.B., Temple University, 1947; Ph.D., Massachusetts Institute of Technology, 1951; Bell Telephone Laboratories, 1951-1956; Hoffman Semiconductor Division of Hoffman Electronics Corporation, 1956-. Between 1949-51 he was a research assistant at the Research Laboratories of Electronics at M.I.T. where he was concerned with cryogenic research. At Bell Telephone Laboratories, Mr. Prince was concerned with the physical properties of semiconductors and semiconductor devices and was associated with the development of silicon devices, including the Bell Solar Battery and the silicon power rectifier.

Mr. Prince is a member of the I.R.E., the American Physical Society, the Association for Applied Solar Energy, the Electrochemical Society and Sigma Xi.

W. W. RIGROD, B.S. in E.E., Cooper Union Institute of Technology, 1934; M.S. in Engineering, Cornell University, 1941; D.E.E., Polytechnic Institute of Brooklyn, 1950; State Electrotechnical Institute, Moscow, U.S.S.R., 1935-39; Westinghouse Electric Corporation, 1941-51; Bell Telephone Laboratories, 1951-. His work has been related principally to the study and development of electron tubes, both the gaseous-discharge and the high-vacuum types. He is a member of the American Physical Society, I.R.E. and Sigma Xi.

STEPHEN O. RICE, B.S., Oregon State College, 1929; California Institute of Technology, Graduate Studies, 1929-30 and 1934-35; Bell Telephone Laboratories, 1930-. In his first years at the Laboratories, Mr. Rice was concerned with the non-linear circuit theory, with special emphasis on methods of computing modulation products. Since 1935 he has served as a consultant on mathematical problems and in investigations of the telephone transmission theory, including noise theory, and applications of electromagnetic theory. Fellow of the I.R.E.

ERLING D. SUNDE, E.E., Technische Hochschule, Darmstadt, Germany, 1926; Brooklyn Edison Company, 1927; American Telephone and Telegraph Company, 1927-1934; Bell Telephone Laboratories, 1934-. Mr. Sunde's work has been centered on theoretical and experimental studies of inductive interference from railway and power systems, lightning protection of the telephone plant, and fundamental transmission studies in connection with the use of pulse modulation systems. Author of *Earth Conduction Effects in Transmission Systems*, a Bell Laboratories Series book. Member of the A.I.E.E., the American Mathematical Society, and the American Association for the Advancement of Science.

HAROLD S. VELORIC, B.A., University of Pennsylvania, 1951; M.A., 1952, Ph.D., 1954, University of Delaware; Bell Telephone Laboratories, 1954-. Between 1951-4 he was a research fellow concerned with the synthesis and analysis of boron and silicon compounds. Since joining the Laboratories, Mr. Veloric has been concerned with the properties and development of solid state devices. He has been associated with the development of several classes of silicon diodes, including power rectifiers, voltage-reference and computer diodes. Dr. Veloric is a member of the American Chemical Society and the Electrochemical Society.

THE BELL SYSTEM

Technical Journal

VOTED TO THE SCIENTIFIC AND ENGINEERING
ASPECTS OF ELECTRICAL COMMUNICATION

VOLUME XXXVI

SEPTEMBER 1957

NUMBER 5

OCT 10 1957
Oceanographic Information for Engineering Submarine Cable
Systems C. H. ELMENDORF AND B. C. HEEZEN 1047

Resistance of Organic Materials and Cable Structures to Marine
Biological Attack L. R. SNOKE 1095

Dynamics and Kinematics of the Laying and Recovery of Sub-
marine Cable E. E. ZAJAC 1129

Theory of Curved Circular Waveguide Containing an Inhomogeneous Dielectric
S. P. MORGAN 1209

Circular Electric Wave Transmission in a Dielectric-Coated Wave-
guide H. G. UNGER 1253

Circular Electric Wave Transmission Through Serpentine Bends
H. G. UNGER 1279

Normal Mode Bends for Circular Electric Waves H. G. UNGER 1292

Bell System Technical Papers Not Published in This Journal 1308

Recent Bell System Monographs 1313

Contributors to This Issue 1317

THE BELL SYSTEM TECHNICAL JOURNAL

ADVISORY BOARD

A. B. GOETZE, *President, Western Electric Company*

M. J. KELLY, *President, Bell Telephone Laboratories*

E. J. MCNEELY, *Executive Vice President, American Telephone and Telegraph Company*

EDITORIAL COMMITTEE

B. MCMILLAN, *Chairman*

S. E. BRILLHART

A. J. BUSCH

L. R. COOK

A. C. DICKIESON

R. L. DIETZOLD

K. E. GOULD

E. I. GREEN

R. K. HONAMAN

H. R. HUNTLEY

F. R. LACK

J. R. PIERCE

EDITORIAL STAFF

W. D. BULLOCH, *Editor*

R. L. SHEPHERD, *Production Editor*

T. N. POPE, *Circulation Manager*

THE BELL SYSTEM TECHNICAL JOURNAL is published six times a year by the American Telephone and Telegraph Company, 195 Broadway, New York 7, N. Y. F. R. Kappel, President; S. Whitney Landon, Secretary; John J. Scanlon, Treasurer. Subscriptions are accepted at \$5.00 per year. Single copies \$1.25 each. Foreign postage is 65 cents per year or 11 cents per copy. Printed in U. S. A.

THE BELL SYSTEM TECHNICAL JOURNAL

VOLUME XXXVI

SEPTEMBER 1957

NUMBER 5

Copyright 1957, American Telephone and Telegraph Company

Oceanographic Information for Engineering Submarine Cable Systems*

By C. H. ELMENDORF and BRUCE C. HEEZEN

(Manuscript received June 4, 1957)

Information on the environment in which submarine cable systems are placed is of vital interest in designing, selecting routes for, placing, and repairing this type of communications facility. Existing data are summarized and evaluated, and their application to submarine cable systems is considered.

I. INTRODUCTION

1.1 General

Oceanography, broadly defined, includes the study of all aspects of the oceans. As a science, it is concerned with gathering data and devising theories which describe and explain the past, present, and the future of the oceans. Oceanography includes physical description of the topography, sediments, and temperature of the ocean bottom; investi-

* This is Lamont Geological Observatory Contribution No. 251. Dr. Heezen is a member of the Lamont Geological Observatory of Columbia University. Additional information on this subject is being presented by Dr. Heezen in a publication by the Geological Society of America.

gation of currents and circulation; study of the geology of the earth's crust under the ocean; and investigation of biological factors.

In designing, in finding the best route for, in laying, and in repairing a submarine cable system one can benefit from as detailed a knowledge of the ocean floor as can be obtained. However, the vastness, complexity, and inaccessibility of the ocean bottom make its study difficult. One must depend on limited data, interpreted with the aid of a knowledge of the earth sciences. The acquisition of specific engineering information is further complicated by the inaccuracies of much of the existing data, and the rudimentary nature of many present theories. Yet, by culling, codifying, interpolating, and interpreting the data gathered during the past hundred years, much can be learned that is applicable to particular cable routes. Further, methods now exist for surveying and describing a route with a thoroughness and accuracy that will permit many refinements in the engineering of future submarine cable systems.

In this paper specific problems of immediate interest in the engineering of submarine cable systems are discussed in order to give a perspective of the use of such data in current applications. Emphasis is placed on the state of existing knowledge and on the accuracy of available data. The work reported forms a foundation for more detailed studies of specific routes and for the application of knowledge which will be derived from rapidly expanding oceanographic studies to the particular problems of submarine cable systems.

1.2 Application To Submarine Cable Systems

How are oceanographic information and technique applied in the selection and description of the detailed path of a new cable? First, variations on a direct route must be examined to avoid ocean bottom conditions which may result in cable failures. Studies of telegraph cable fault records indicate that many of the deep sea cable breaks occur where cables pass over sea mounts, canyons and areas susceptible to turbidity currents, and an effort must be made to avoid such hazards. Topographic studies form the basis for both initial route selection planning and for a preliminary description of the selected route.

This description will include, where data are available, an exaggerated depth profile uncorrected for angle of the sound beam of the sonic depth recorder, a corrected 1:1 depth profile where necessary, and a temperature profile. Also, bottom characteristics, including photographs, can be collected for the particular route. These data will be essential in planning a detailed survey of the route, and they will provide

preliminary information for the determination of the amount of cable required, for system transmission planning and for studies of cable-laying techniques.

One of the fundamental requirements in cable laying is the deposit of a sufficient amount of cable to cover irregularities of the bottom without introducing dangerous suspensions and without laying a wasteful amount of excess slack. Satisfaction of this requirement will require the most detailed possible knowledge of bottom topography, coupled with knowledge of the kinematics of the cable laying process.¹

A determination of required cable strength does not directly call for an extremely accurate knowledge of bottom depth and contour. The required cable strength is determined by cable tension during recovery, which is two or more times that experienced during laying. Although the required strength is directly proportional to the depth, it is also affected to a major degree by the ship speed, cable angle during recovery, and the ship motion caused by waves. The ship motion can be controlled to some extent by seamanship and choice of the time at which the recovery is to be made.

Further, the design strength of the cable is in large measure determined by the strength of available steel and the amount of steel that can be accommodated in an economical over-all design. Thus, uncertainties of 5 to 10 per cent in the maximum depth of the water on a route would not affect the cable design. Yet, during a critical recovery situation, a more accurate knowledge of depth would be useful in planning and executing the operation.

The integrity of the cable will depend on the choice of materials to withstand biological factors and wear on the ocean bottom. A companion paper² discusses the resistance of likely cable-sheath materials to attack by marine borers and bacteria.

II. TOPOGRAPHY

2.1 *General*

Mapping the ocean bottom involves depth measurements coupled with a knowledge of the location on the earth's surface of the points at which these measurements are made. There exists a vast store of data on the depths of the oceans taken by hundreds of observers and expeditions over the past hundred years. It is not surprising that many of these data are of questionable accuracy due to errors in navigation, soundings, plotting, and interpretation. Thus, before existing data can be used for the engineering of submarine cable systems, it is essential

TABLE I — MARINE NAVIGATION SYSTEMS

Name	Description	Range	Accuracy	Coverage
LORAN	Measures time difference between pulses transmitted by two or more shore station groups.	900 miles*	$\pm \frac{1}{2}$ mile at short range	Good over N. Atlantic steamer routes; spotty elsewhere.
SHORAN	Radar system: shipboard equipment triggers two shore stations to determine position.	Line of sight	± 50 feet	Nil
DECCA NAVIGATOR	Measures phase difference in continuous waves transmitted by two or more shore stations.	240 miles	± 100 yards	Good around British Isles practically nil elsewhere.
LORAC	Measures phase difference in continuous waves transmitted by three shore stations.	100 miles	± 50 yards	Used as a temporary installation for special purposes.
CONSOL	Radio determination of great circle bearing of transmitting station.	1500 miles	$\pm 0.2^\circ$ to 1°	Eastern North Atlantic
SOFAR	Experimental underwater sound system; shore stations determine distance to under-water explosion.	4000 miles†	± 4 miles	Eastern Pacific only
E. P. I.	An electronic position indicator which measures distances from ship to each of two shore stations, using a pulse technique.	350 to 500 miles (max.)	± 75 yards (in measured distances)	Developed by U.S.C. & G.S. for surveying purposes, temporary installations.
TACAN	Provides range and bearing from transmitter.	Line of sight	Range $\pm \frac{1}{2}$ mile, bearing $\pm \frac{3}{4}^\circ$	Experimental only

* Ground wave; Sky wave range up to 1400 miles with decreased accuracy (15 miles at max. range).

† Range limited by size of basin and topographic obstructions.

that it be compiled, evaluated, and culled with due regard for the sources, inherent errors, and possible misinterpretations. The data remaining after this type of review can then be combined with other knowledge from the earth sciences to provide the best available maps and profiles.

2.2 *Instrumentation for Topographic Studies*

2.2.1 *Navigation*

Celestial navigation, depending on hand sextant sights of heavenly bodies and subsequent solution of an astronomical triangle, is the only method available in all parts of the world. It is, of course, useless during periods of poor visibility. Positional accuracy of $\pm\frac{1}{2}$ mile can be obtained, although the usual accuracy is ± 2 to 5 miles. During the interval between sights, an estimated position is computed by dead-reckoning methods, assuming the ship's speed and course to be known. Speed determined by propeller revolutions or by a pitometer log is subject to the influence of winds, seas, and currents. The ship's heading is indicated by compass, but the course actually traveled by the vessel is affected by all the same variables that affect the speed. Starting from a celestial or radio fix, a good navigator can plot a dead-reckoning track with little error over short periods of time under ideal conditions. The probable error increases rapidly with elapsed time under less than ideal conditions.

Some of the existing radio navigational systems and their limitations are listed in Table I.

2.2.2 *Echo Sounding Errors and Corrections*

Present-day echo sounders, operating with a total beam angle of about 60 degrees, indicate the distance to the best reflectors within the effective cone of sound. Neglecting scattering layers, which can usually be eliminated from consideration by interpretation, the best reflector will coincide with the bottom immediately below the transducer if the bottom is horizontal or if the highest point on the bottom is immediately below the transducer. Under any other condition the echo sounder indicates less than the true depth. Where the bottom is very uneven or rocky, a multiplicity of echoes are returned and recorded. These considerations necessitate careful evaluation of the data and its correction for slope.

Since the echo sounder actually measures the time interval between transmission of a pulse and receipt of an echo, the timing mechanism

is the heart of the indicating and recording equipment. However, many types of echo sounders have poor timing mechanisms which depend upon mechanical governors, poorly regulated AC power supplies, and friction drives. Recently, a Precision Depth Recorder (PDR) having an instrument accuracy of better than one fathom in 3,000 has been developed.³ The equipment is used with standard deep-sea sounding equipment. This PDR will perform the timing function with an accuracy better than one part in a million.

After obtaining a record of the time between the outgoing and received pulses the data must be converted to depth, employing velocity corrections and slope corrections. True sound velocity is obtained by

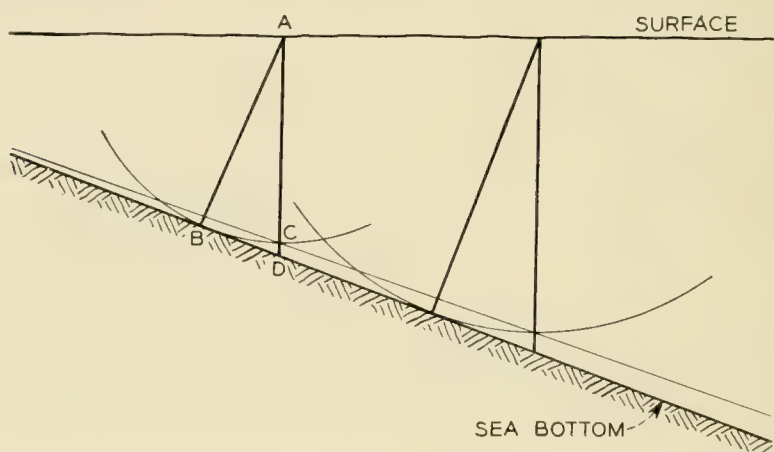


Fig. 1 — Slope correction of echo soundings.

reference to tables, or by computation based on simultaneous sea water temperature and salinity measurements. The distance traveled by the echo, obtained from the velocity-corrected data, is converted to depth by the slope correction.

The PDR sends out a ping and records an echo once each second. If the sounding vessel is under way at 10 knots, the individual soundings are about 17 feet apart. When the outgoing ping exactly coincides with the returning echo, a gating circuit reduces the frequency of soundings to once in 2 seconds corresponding to a spacing of about 34 feet at 10 knots.

Fig. 1 illustrates the problem of slope correction on an ideal slope, where the ship is steaming at right angles to the slope. At sounding position A, the echo is returned from B, the nearest point on the bottom. On the uncorrected profile, this would be interpreted as a vertical sounding AC. To obtain the true vertical depth at A, an amount CD must

be added to the depth read from the echo trace. This is done graphically by swinging arcs representing the echo distances, using the distance between soundings as the arc center spacings. The envelope of a succession of such arcs is the best approximation to the actual bottom configuration that can be made. This method, of course, requires that the sounding track be run approximately at right angles to the slope, and thus, that the trend of the topography be determined. In addition, when the relief becomes more complicated, constructing and interpreting the envelope becomes much more difficult. Although all hills are shown on echograms, small valleys are often completely missed because their width is much less than the breadth of the cone of sound. This problem is partly eliminated by the use of a PDR, where second echoes indicate the existence of valleys not recognized on standard echo sounders.

Fig. 2 (a) shows the trace made by a PDR in passing over a rugged slope in mid-Atlantic. The multiple echo on the left hand side of the figure illustrates how echoes from the wide sound cone are returned from different parts of a steep slope. Similarly, the multiple echoes in the center show a deep valley with energy from the same pulse reflected from the steep slopes as well as the bottom. Fig. 2 (b) shows a corrected profile constructed from the echogram illustrated in Fig. 2 (a), and Fig. 2 (c) is a profile of the same track with 40:1 vertical exaggeration, which is typically used in describing bottom topography.

To construct a profile and interpret each sounding taken would require a prohibitive amount of work. Thus, in preparing contours and uncorrected profiles for ordinary mapping purposes, a sufficient number of the initial echoes are taken to allow a reasonably accurate profile to be constructed. Where 1:1 corrected profiles are required, all the echoes are used and the best possible envelope representing the bottom is constructed.

2.3 CHARACTERISTICS OF AVAILABLE INFORMATION

2.3.1 *Sources of Data*

The vast majority of data for a study of submarine topography are obtained as a series of "soundings," or depths below sea level.* Prior to the development of echo sounding apparatus, soundings were made by measuring the length of either a weighted piano wire or hemp line which was lowered to the bottom. This method of line or wire sounding

* Features to be studied by detailed soundings can often be located by measurements of the earth's magnetic and gravitational fields and by acoustic methods.

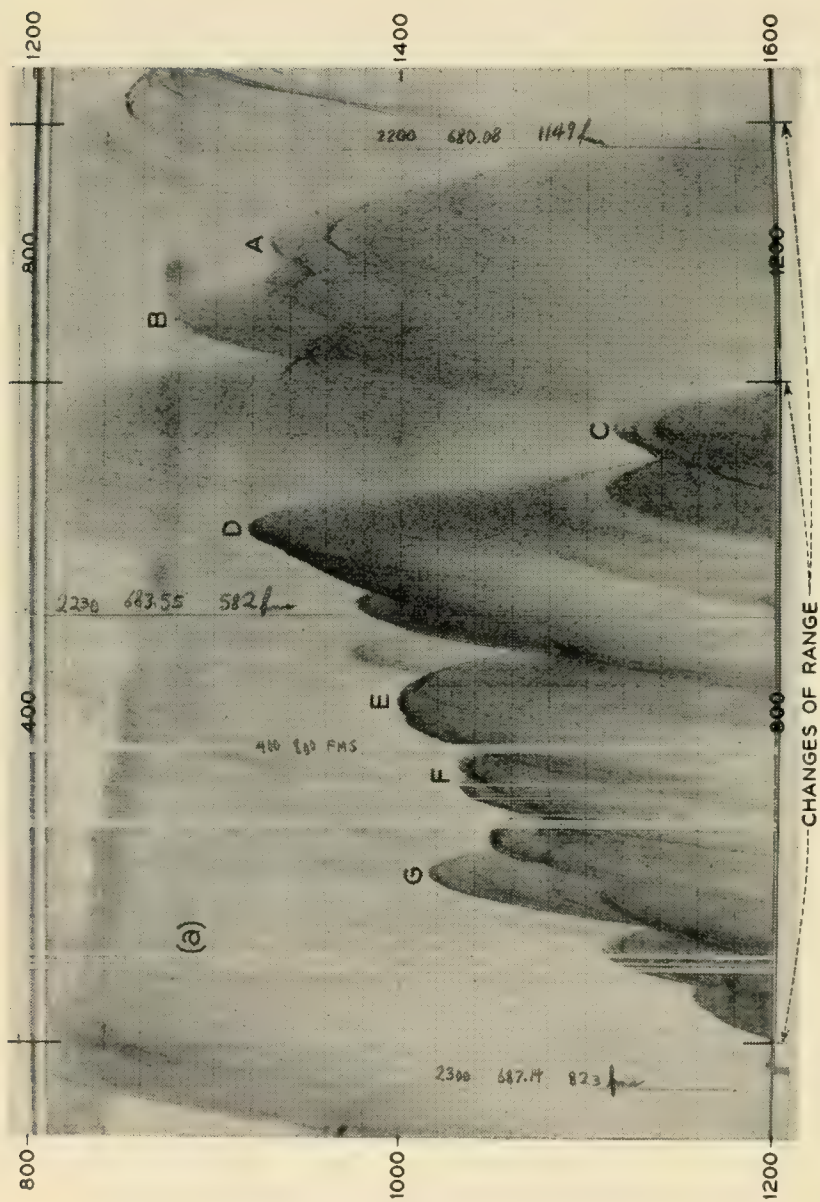
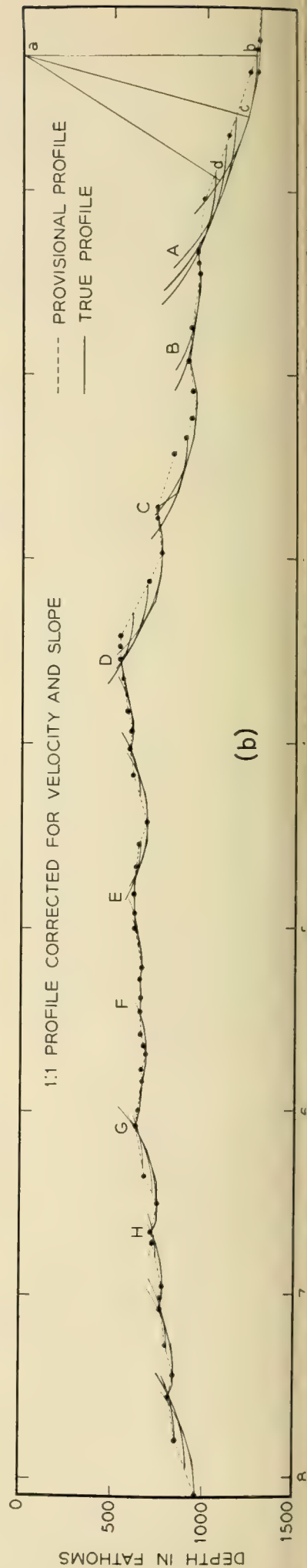


Fig. 2 — (a) Precision-Depth-Recorder record of a peak in the Mid-Atlantic Ridge. Record shows multiple 400-fathom scales. On original record 400-fathom range is represented by 18 inches and one hour or about 10 miles by 24 inches. Light horizontal lines are at 20-fathom intervals. (b) Same profile with no vertical exaggeration, with corrections for slope. (c) Profile at 40:1 exaggeration plotted from first echoes shown in (a).



was fraught with error because there was no adequate sensing device to indicate when the bottom was reached, and because the sounding line might be swept far from a true vertical by the action of currents between the surface and the bottom. These difficulties could cause errors as great as 50 per cent. Echo sounding was thus a tremendous improvement despite its own inherent inaccuracies.

Ocean bottom soundings are available from a number of sources including, in this country, government agencies such as the Navy's Hydrographic Office and the U. S. Coast and Geodetic Survey, and private oceanographic institutions. Abroad, national hydrographic offices, such as the British Admiralty's Hydrographic Department, the Japanese Hydrographic Office, and the International Hydrographic Bureau, collect and publish soundings.

Depending upon the organization which compiles the soundings, various corrections are applied to the raw data, each organization selecting both the corrections it wishes to apply and the method of application. If all soundings for an area off the continental shelf of the United States were compiled, there might be available soundings in feet, corrected for velocity, supplied by the U. S. Coast and Geodetic Survey; uncorrected soundings in fathoms supplied by the U. S. Navy Hydrographic Office; similar soundings supplied by Lamont Geological Observatory; corrected soundings in fathoms supplied by the British Admiralty; and corrected soundings in meters supplied by the International Hydrographic Bureau. In addition, there would be a quantity of hemp, wire and discrete echo soundings on published charts. One difficulty in using the soundings printed on the published charts arises from the fact that hemp line, wire, and echo soundings of all types, both corrected and uncorrected, are all plotted on the same chart usually without designation as to method or corrections.

Table II summarizes the methods of presenting sounding data used by various agencies.

Several methods of recording continuous depth records are used. The most common, and least satisfactory, is to read the echo sounder and plot the sounding at the appropriate point on the chart at discrete intervals, say every 10 or 20 minutes. This achieves an orderliness in printing but has the disadvantage that canyons, mountains, or other features which cannot be adequately represented by such spacing are ignored and obscured in the plotted soundings. Better results are obtained by use of the "texture method," where soundings are recorded at each crest, valley, or change of slope, and soundings at uniform time intervals are only used in areas where a continuous slope extends for

TABLE II — METHODS OF PRESENTING SOUNDING DATA

Source	Depth Units	Sounding Velocity	Velocity Correction	Slope Correction
Coast and Geodetic Survey, U. S. Dept. of Commerce.....	Fathoms	820 or 800 fm/sec	Note A	Note B
Hydrographic Office, U. S. Navy...	Fathoms	800 fm/sec	No	No
Lamont Geological Observatory.....	Fathoms	800 fm/sec	No — Note E	No — Note E
British Admiralty Hydrographic Dept.....	Fathoms	820 fm/sec	Yes — Note C	No
International Hydrographic Bureau.....	Meters	Note D	Yes	Note D
Japanese Hydrographic Office.....	Meters	1500 meters/sec	?	?

1. All agencies use Mercator projection charts.
2. All deep sea soundings on 4" = 1° longitude charts; various larger scales used near shore.
- A. USC & GS usually makes velocity correction according to data taken at time of sounding.
- B. USC & GS makes slope and drift corrections where deemed necessary.
- C. Admiralty data velocity corrections are made according to D. J. Matthews.⁴
- D. International Hydrographic Bureau takes no data of its own, but publishes data received from various surveyors.
- E. Although corrections are made in surveys of specific areas, soundings are first compiled in uncorrected form.

many miles. This produces an uneven spacing of numbers on the chart, and in some areas the soundings have to be crowded together to show all crests and valleys. In this method the sounding is written alongside a dot which represents the location of the sounding. Another method is to write the sounding without a dot but centered over the place where the sounding was taken. This produces a more pleasing drawing but all detail must be left out in complicated areas since the physical size of the letters limits the number of soundings that can be so recorded on the chart.

2.3.2 *Methods of Evaluation*

Evaluation of topographic data starts with a comparison of data from all the different sources, with all the soundings plotted to the same scale on one set of charts. Soundings from different sources must be reduced to a common base.

When the soundings from all sources are compiled on the same sheet, many obvious discrepancies will be noted. A great number of these

cases can be traced to either gross mistakes in plotting or poor accuracy in navigation which may cause errors in position of up to 25 miles. Usually these gross errors are fairly easy to spot. For example, a large shoal area was shown for many years in deep water east of Georgia, but it now appears that this particular shoal resulted from the misplotting, by one-half degree of longitude, of a series of continental shelf soundings.

Lines of soundings from different sources should indicate the same depth at their intersections, providing a check on the reliability of both or an indication that one or both is suspect. Where there is lack of agreement the sounding and navigational methods should be checked in an effort to find a basis for choosing one set of soundings rather than the other. However, without special knowledge about the individual sounding lines, it is often impossible to decide which line is correct. Uncertainties of 2 to 25 miles in position combined with possible errors in depth determination of 10 per cent and possible gross errors in plotting are indicative of the difficulties that may be encountered.

The extent of the coverage of the Atlantic Ocean with precision sounding tracks is shown on Fig. 3. Sounding tracks taken by the Lamont staff prior to the availability of the PDR are shown on Fig. 4.

2.3.3 *Methods of Presentation*

Relief is usually indicated on maps and charts by any one of a number of devices, including contour lines, profile views, and physiographic sketches. In contouring, after surveying a number of control points and obtaining their exact elevations, the land surveyor sketches in the contour lines between control points while standing on a vantage point such that he can actually see the terrain. In contrast, the oceanographer must sketch contour lines by applying his own interpretation of the submarine processes responsible for the relief in the areas between soundings. The accuracy of contour is, of course, determined by both the number, spacing, and accuracy of the soundings, and the skill and knowledge of the oceanographer.

The International Hydrographic Bureau in Monaco publishes a colored contour chart for the entire world on a scale of 1:10 million, individual sheets of which are revised and republished at 10–25 year intervals.

Profiles, or elevation views along particular tracks, provide a detailed outline of the bottom. The usual practice is to construct exaggerated profiles (40:1 or greater vertical to horizontal scale ratio), such as those illustrated in Figs. 2(c) and 5. Where accuracy of the highest order



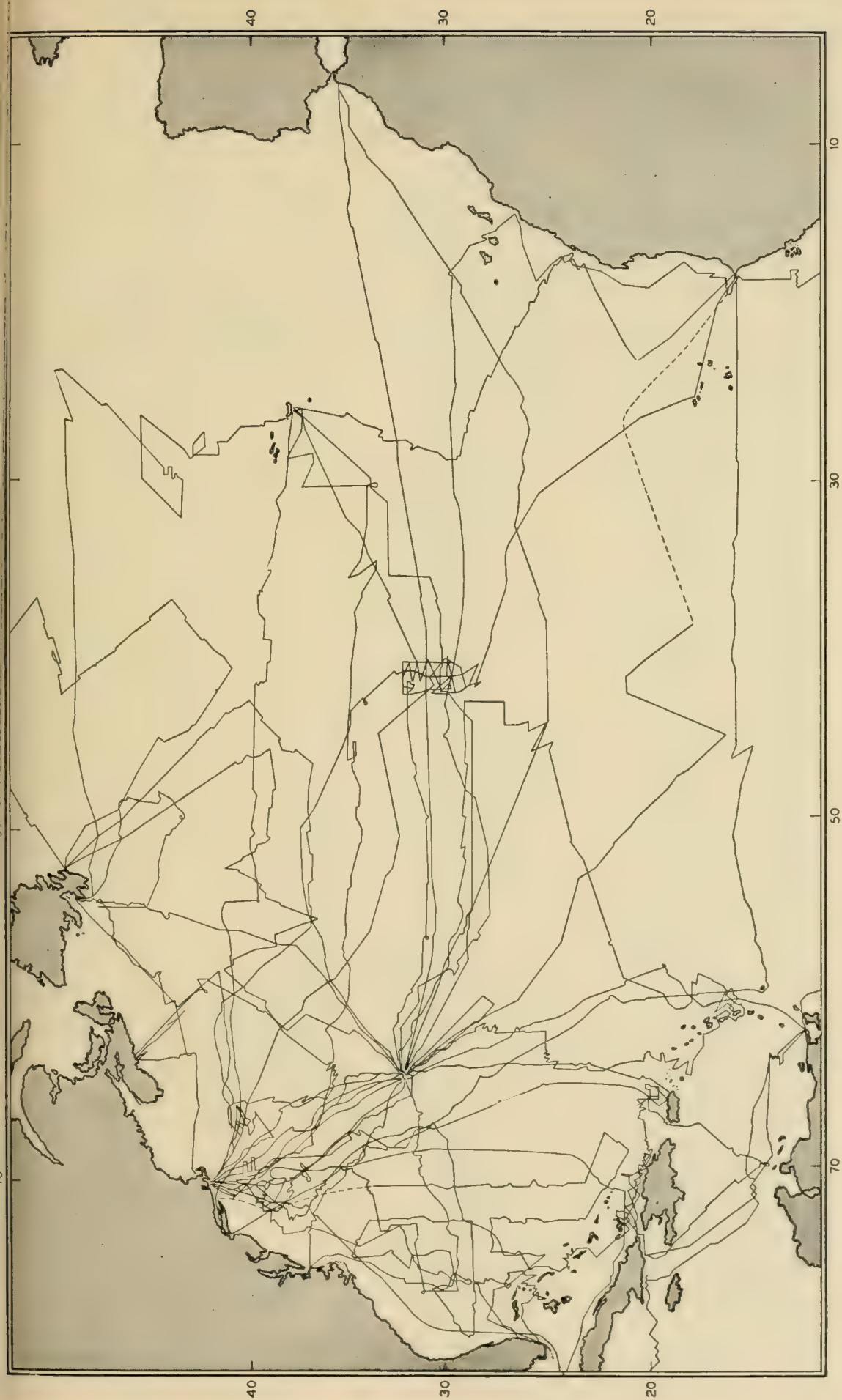


Fig. 4 — Tracks of research vessels employing fairly good but nonprecision depth recorders. Tracks mainly of *R. V. Atlantis*, 1946–1952.

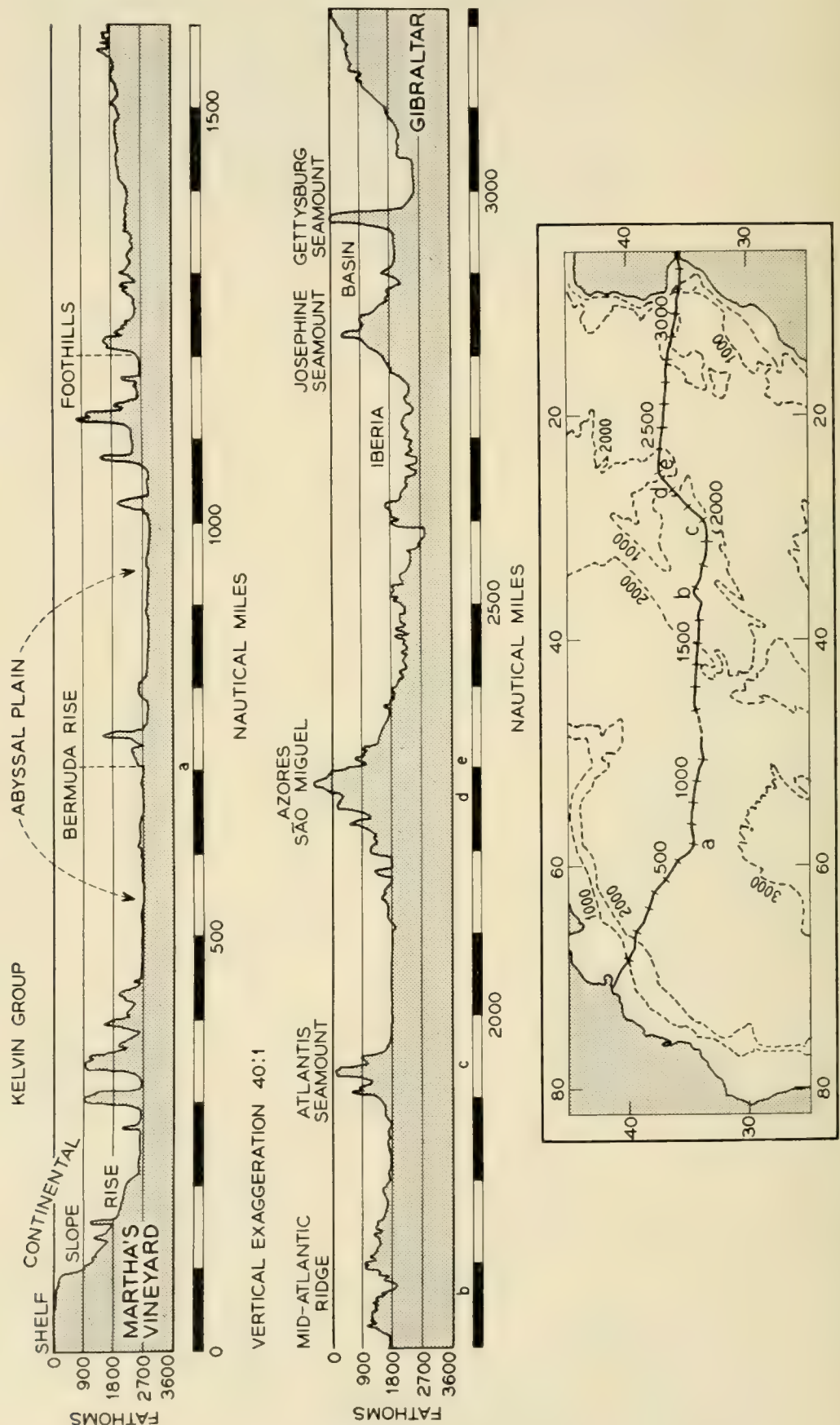


Fig. 5 -- Transatlantic topographic profile, Massachusetts to Gibraltar. Vertical exaggeration 40:1.

is required, 1:1 profiles as illustrated in Fig. 2(b) are prepared by careful interpretation of precision depth records and their correction for actual sound velocity and slope.

The physiographic diagram (in envelope inside rear cover) was prepared by combining evaluated and corrected depth measurements with all other ocean bottom and geological information. The construction of these diagrams is preceded by the preparation of 40:1 exaggerated profiles along all available sounding tracks through the region under consideration. After study and evaluation of the profiles, areas of different texture are defined, and the various physiographic provinces are outlined. The relief shown on the profiles is then drawn in perspective view along the appropriate track of the chart. After all available tracks have been drawn, the remaining blank areas are sketched in, basing the interpretation on geological information.

Fig. 6 shows the physiographic provinces of the North Atlantic and the existing submarine cables on a great circle chart. A plot such as this is useful in preliminary studies of new routes. More detailed charts to a scale of 1:1 million showing contours, actual sounding tracks, existing cable routes and cable fault records can be prepared for specific engineering of new cable routes.

2.4 *North Atlantic Topography*

The relief of the continents naturally divides itself into mountain ranges, plateaus, and plains — physiographic provinces which can be recognized by distinct differences in topography, form and texture. Geologists recognize these differences as directly related to underlying geological structures. The topography of the ocean floor can also be divided into physiographic provinces similar to those familiar on land. In general, the relief of the deep sea topography is greater than on land due to the fact that the smoothing effects of erosion are less under the deep sea. Fig. 7 shows the relief along a continuous line extending from Peru across South America and thence across the South Atlantic.

The three major divisions of the Atlantic Ocean — Continental Margins, Ocean Basins, and Mid-Atlantic Ridge — each occupy about one third of the ocean floor. Since detailed topographic information is available for so little of the area covered by the oceans, the following descriptions of the relief must of necessity be general. Where specific details are available, they are presented as examples. The area to be described, the North Atlantic, is well outlined on the Physiographic Diagram. Frequent reference to this illustration should help to make the word descriptions more meaningful.

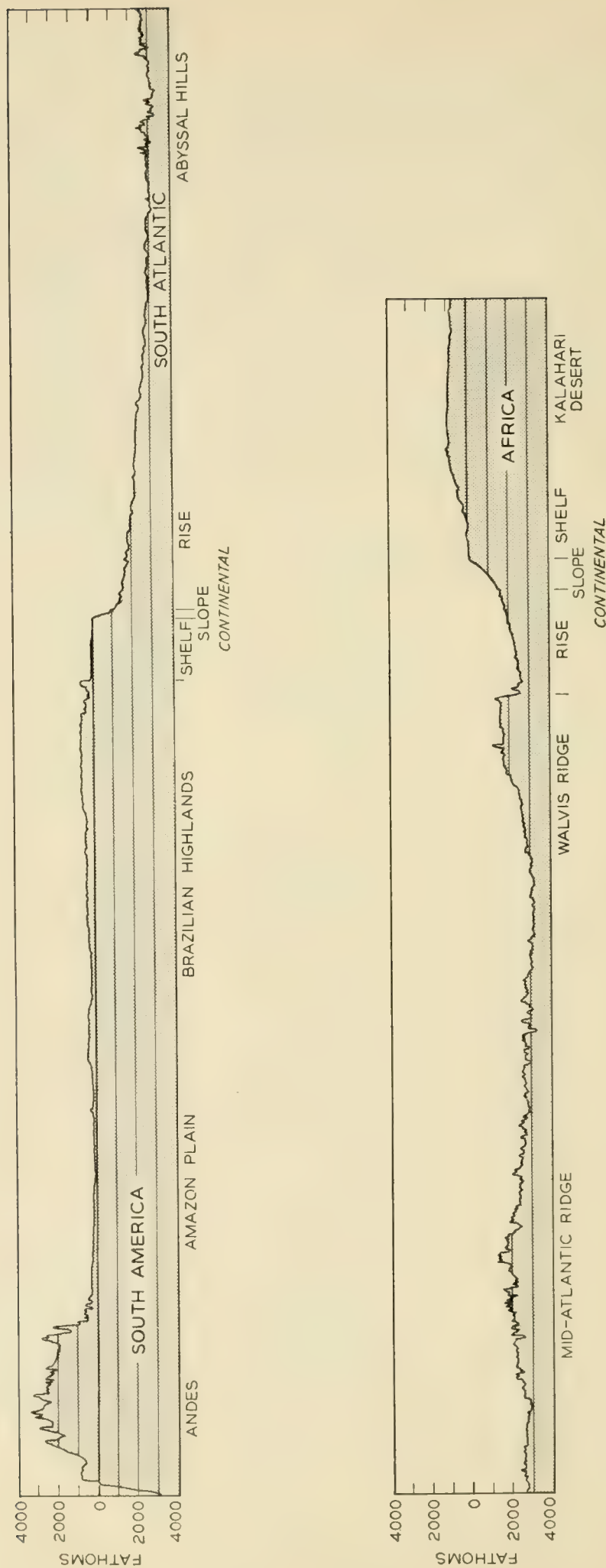


Fig. 7 — Topographic profile from the Pacific across South America and the South Atlantic to Africa. 40:1 exaggeration.

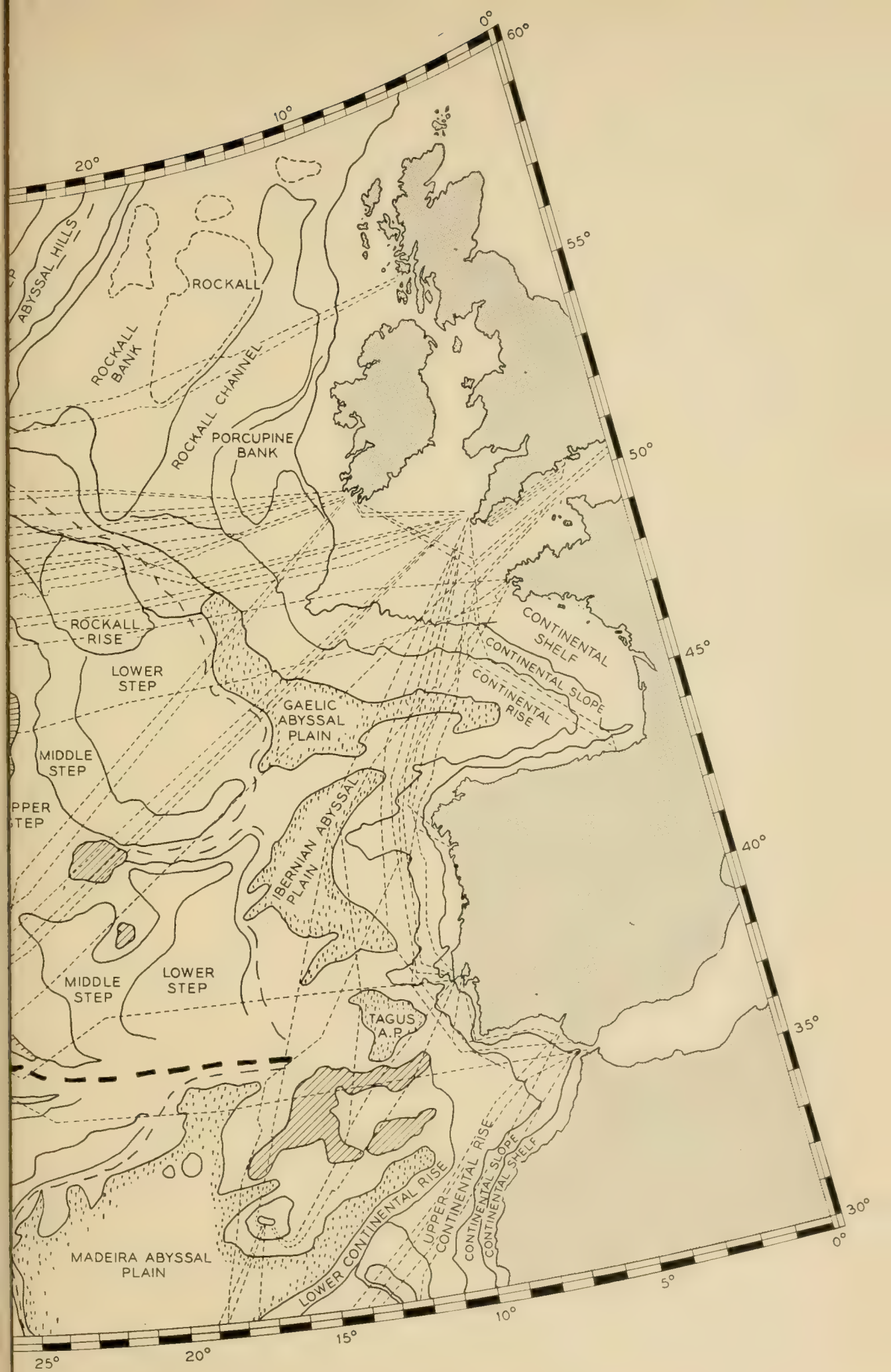




Fig. 6 — Great-circle chart showing physiographic provinces and existing submarine cable routes (dotted lines). (On this chart any straight line represents a great circle.)

2.4.1 *Continental Margins*

The first area considered is the continental shelf which may be characterized as the shallow (0–100 fathoms) submarine terrace bordering the continents, extending seaward hundreds of miles in some localities. The shelf terminates where the bottom gradient increases suddenly from the shelf average of approximately 0.1° to the continental slope average of more than 4° . This change occurs in depths ranging from 20 to more than 100 fathoms. The shelf is a continuation of the coastal plain and displays the same numerous small irregularities as the plain. For example, off the United States east coast — from Cape Cod south — the shelf is relatively flat, but with many small (10-fathom) hills and ridges. This shelf varies in width from about 150 miles off Cape Cod to its virtual disappearance off the east coast of Florida. This region of smooth shelf terminates in depths of about 60 fathoms between Cape Cod and Cape Hatteras, but in only 20–30 fathoms south of Hatteras. It is crossed by at least three submerged valleys, off the Hudson River and Delaware and Chesapeake Bays.

North of Cape Cod the shelf presents a somewhat different pattern with many basins and troughs pitting the shelf whose width increases to 240 miles off Newfoundland. Characteristic of this region are the extensive off-shore shoals (the “Banks”), with depths of about 30 fathoms, which are found to seaward of areas with depths up to 100 fathoms.

Outside the continental shelf the bottom drops comparatively rapidly down the continental slope. The top of the continental slope usually lies near the 100-fathom contour but the base lies in depths varying from 700–2,000 fathoms, depending on the area. The typical slope has a gradient of approximately 1:13 (about $4\frac{1}{4}^\circ$). The base is marked by a sharp change in the seaward gradient, from values greater than 1:50 (1°) to values less than 1:200 ($\frac{1}{4}^\circ$). The continental slope is dissected by numerous submarine canyons, comparable in dimensions to the canyons of mountain slopes. In some cases these canyons connect with shelf channels but generally just slightly indent the shelf edge. Some of the canyons continue across the continental rise (at the foot of the slope), but most of them apparently flatten out and disappear in the rise.

The continental rise located at the foot of the continental slope is a gently sloping apron with gradients varying from about 1:200 ($0^\circ 17'$) to nearly 1:1,000 ($0^\circ 3\frac{1}{2}'$). Minor relief features are rare although there are submarine canyons and occasional protruding seamounts. The submarine canyons having steep, V-shaped walls which may approach the

2.4.1 *Continental Margins*

The first area considered is the continental shelf which may be characterized as the shallow (0–100 fathoms) submarine terrace bordering the continents, extending seaward hundreds of miles in some localities. The shelf terminates where the bottom gradient increases suddenly from the shelf average of approximately 0.1° to the continental slope average of more than 4° . This change occurs in depths ranging from 20 to more than 100 fathoms. The shelf is a continuation of the coastal plain and displays the same numerous small irregularities as the plain. For example, off the United States east coast — from Cape Cod south — the shelf is relatively flat, but with many small (10-fathom) hills and ridges. This shelf varies in width from about 150 miles off Cape Cod to its virtual disappearance off the east coast of Florida. This region of smooth shelf terminates in depths of about 60 fathoms between Cape Cod and Cape Hatteras, but in only 20–30 fathoms south of Hatteras. It is crossed by at least three submerged valleys, off the Hudson River and Delaware and Chesapeake Bays.

North of Cape Cod the shelf presents a somewhat different pattern with many basins and troughs pitting the shelf whose width increases to 240 miles off Newfoundland. Characteristic of this region are the extensive off-shore shoals (the “Banks”), with depths of about 30 fathoms, which are found to seaward of areas with depths up to 100 fathoms.

Outside the continental shelf the bottom drops comparatively rapidly down the continental slope. The top of the continental slope usually lies near the 100-fathom contour but the base lies in depths varying from 700–2,000 fathoms, depending on the area. The typical slope has a gradient of approximately 1:13 (about $4\frac{1}{4}^\circ$). The base is marked by a sharp change in the seaward gradient, from values greater than 1:50 (1°) to values less than 1:200 ($\frac{1}{4}^\circ$). The continental slope is dissected by numerous submarine canyons, comparable in dimensions to the canyons of mountain slopes. In some cases these canyons connect with shelf channels but generally just slightly indent the shelf edge. Some of the canyons continue across the continental rise (at the foot of the slope), but most of them apparently flatten out and disappear in the rise.

The continental rise located at the foot of the continental slope is a gently sloping apron with gradients varying from about 1:200 ($0^\circ 17'$) to nearly 1:1,000 ($0^\circ 3\frac{1}{2}'$). Minor relief features are rare although there are submarine canyons and occasional protruding seamounts. The submarine canyons having steep, V-shaped walls which may approach the

vertical generally resemble land canyons cut in the sides of mountain ranges. These canyons, some with tributaries, usually follow gently curving to straight courses down the slope from their origins on the shelf and gradually disappear on the ocean floor. Most exhibit sediment filling on their floors, and many have rocky walls. In some areas the base of the continental rise is marked by a range of hills 100 fathoms or less in height and one mile wide or less at the base.

2.4.2 *Ocean Basins*

The continental rise gives way to the abyssal plains which occupy a sizable proportion of the ocean basins. The smooth, nearly flat topography of the abyssal plains was apparently produced by deposition of sands and silts which were carried by turbidity currents from the continental margins via the submarine canyons.

Sea mounts which rise from the abyssal plain have an appearance of being partially buried. The small sea mounts and hills which protrude from the abyssal plain increase in number toward the seaward limit of the plain. The seaward extremity of the abyssal plains frequently occurs where the small hills become so numerous that they occupy the entire area. The margin of the abyssal plain along certain positive features such as the east or west margin of the Bermuda Rise is marked by a sharp rise of the sea floor where the depositional floor has built up against the topographic rise. The larger sea mounts that occur scattered through the abyssal plain also show the same partially buried appearance.

Besides the two lines of large sea mounts which parallel the Mid-Atlantic Ridge on the east and west, there is another major trend of sea mounts, the Kelvin Group, running southeast from New England.

The eastern and western basins of the Atlantic are quite similar but there are several significant differences. The European continental slopes are in general higher, steeper, and more rugged and irregular than those off the North American coast. The continental rise is often absent or poorly developed; in some areas the continental slope descends almost directly to the abyssal plain. In the area north of the Azores the abyssal plains are less well developed than on the west side of the ridge. Rockall, Bill Bailey's, and Lousy Banks, rocky spines running south from the Iceland-Faeroe Ridge, represent features which have no direct analogies in the western basin. In the area between the Azores and Gibraltar numerous sea mounts of large size are encountered more frequently than in the western basin. The northwest margin of Africa bears the closest similarity to the northeast coast of the United States, with an extensive abyssal plain and a well-developed continental rise.

A mid-ocean canyon three miles wide with precipitous walls which drop fifty fathoms to the canyon floor runs down the length of the Labrador and Newfoundland Basins as shown on Fig. 8. Profiles across this canyon are shown in Fig. 9. Other mid-ocean canyons have recently been discovered in the equatorial Atlantic as well as in the basin south of Nova Scotia. Studies of sediments obtained in these canyons suggest that they were cut by turbidity currents.

2.4.3 *Mid-Atlantic Ridge*

The principal topographic feature of the Atlantic is the Mid-Atlantic Ridge which runs the entire length of the Atlantic and continues into the Indian and Arctic Oceans. The ridge is about 1,200 miles wide and can be thought of as a broad swell or arch with varied and generally extremely irregular topography. Along the axis of the ridge is a narrow crest about 60 miles wide with a characteristic median depression which cleaves the crest zone. Depths in the median depression exceed the maximum depths of the adjacent flanks of the ridge out to 100 miles or more. In some cases they reach depths equal to those of the abyssal plains. The tops of the highest peaks of the ridge, excluding those which emerge as islands, lie at about 800 fathoms while the median rift falls to depths of about 2,000 fathoms and locally to depths as great as 2,800 fathoms.

Near the outer margins of the ridge there is a discontinuous line of sea mounts which rise as isolated peaks. The major part of the ridge lies at depths intermediate between the abyssal plains and the central highlands, with extensive areas of flat intermountain basins, particularly in the area just south of the Azores. The earthquake epicenter belt accurately follows the median rift throughout the length of the ridge.

Other areas of the Atlantic which have topography similar to the Mid-Atlantic Ridge include an oval area trending northeast-southwest from Bermuda with a long axis of about 800 miles and a short axis of about 500 miles. The area is characterized, in part, by low irregular relief, but with a number of large sea mounts.

III. NATURE OF THE SEA BOTTOM

3.1 *Methods of Investigation*

There are four principal methods of investigating the nature of the bottom:

- (a) Visual inspections and photographs.
- (b) Physical sampling.

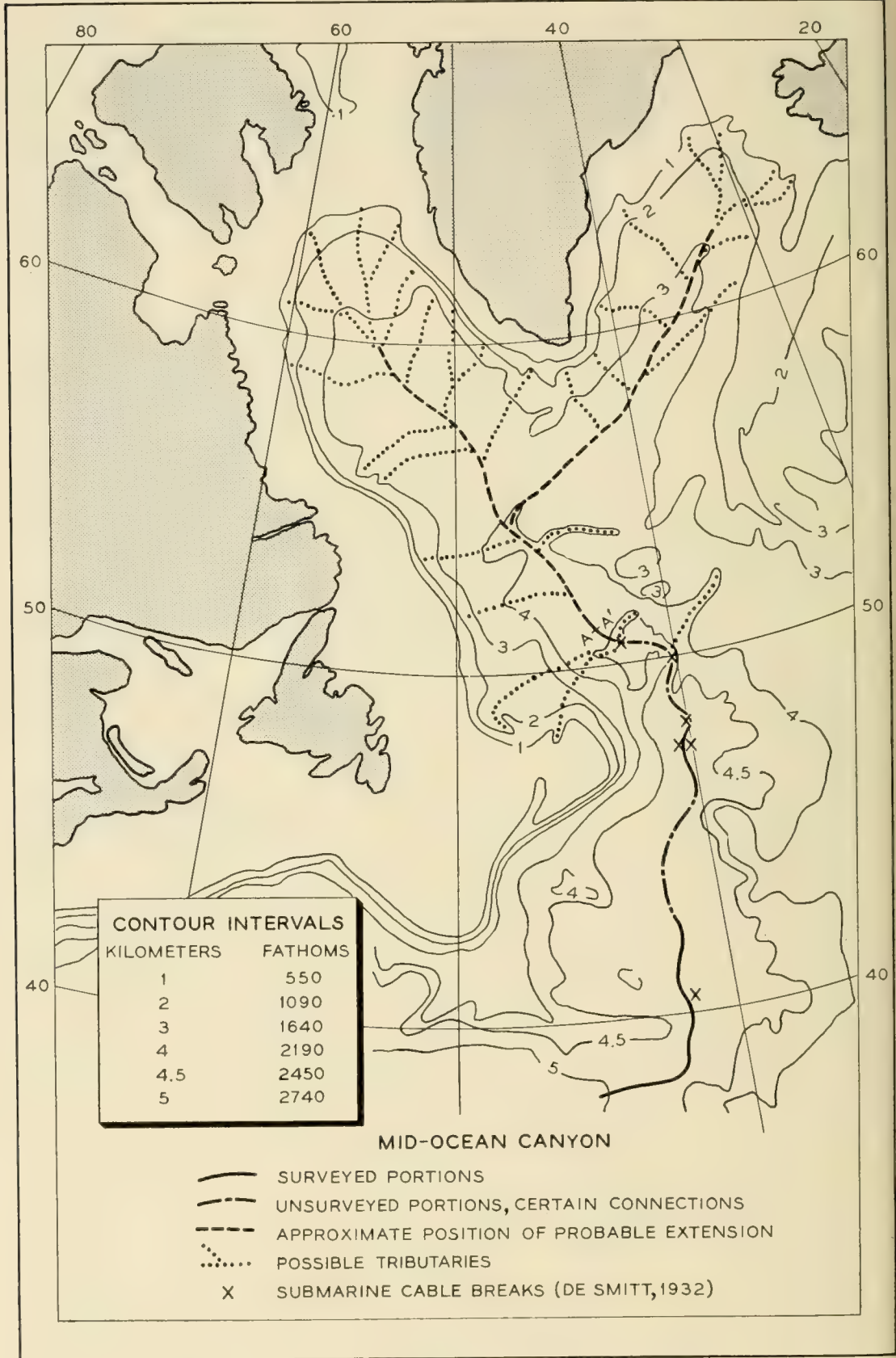


Fig. 8 — Mid-ocean canyon of the Northwest Atlantic.

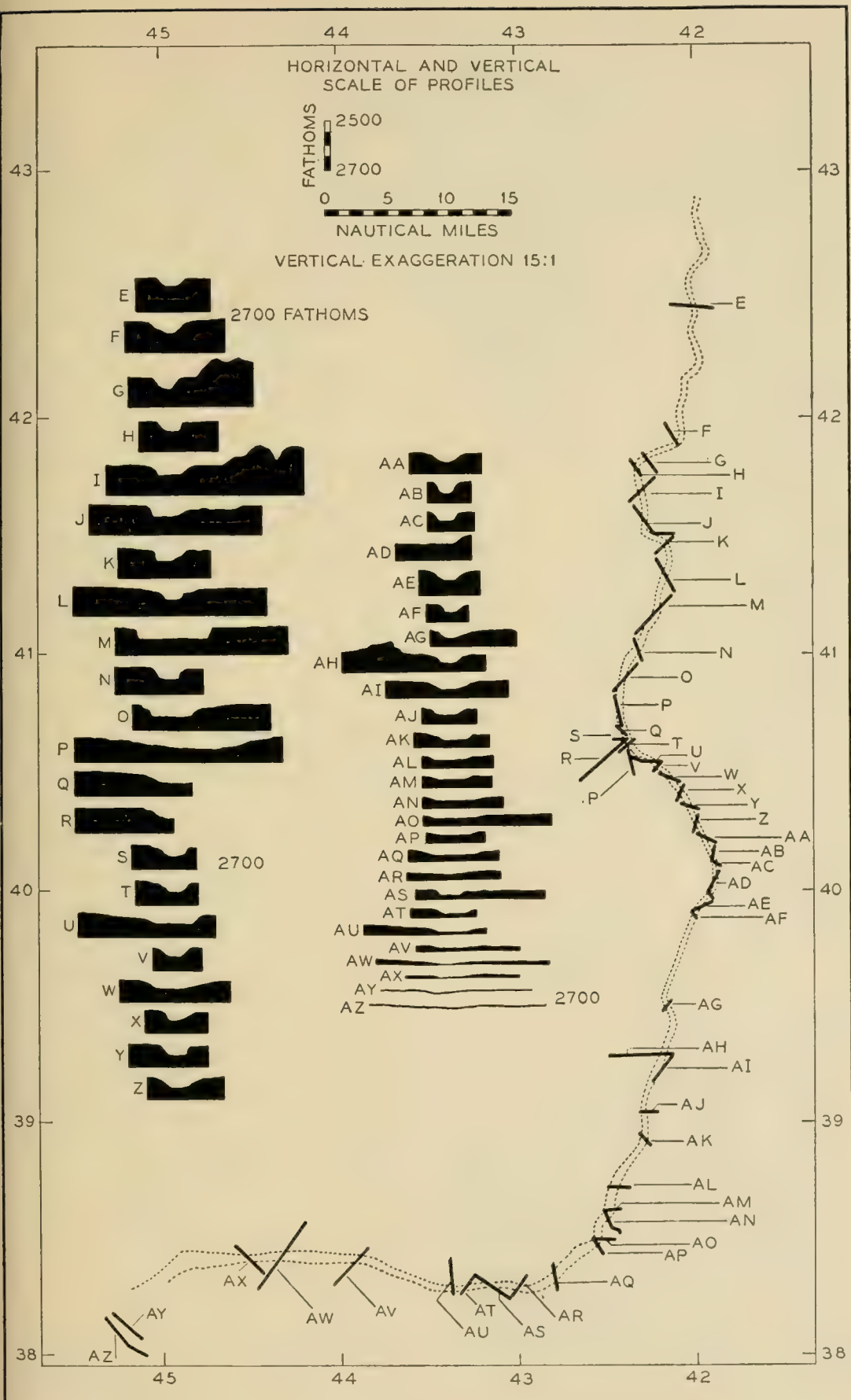


Fig. 9 — Profiles across the mid-ocean canyon.



Fig. 10 — Location of cores, rock dredges, and bottom-photograph sites of Columbia University expeditions in the North Atlantic, 1946-1956.

(c) Sound investigations.

(d) Investigations of magnetic and gravity fields, and heat flow.

Visual inspections have been performed by divers in depths up to only about two hundred feet. Vessels such as the bathyscaphe must be employed for greater depths but their design precludes any extensive observation of the sea floor.

Bottom photographs in shallow waters are fairly numerous but those taken in depths greater than 1,000 fathoms are still rare. Good photographs usually show an area approximately 5 ft by 8 ft and are focussed well enough to make objects and organisms $\frac{1}{8}$ inch in diameter clearly identifiable. A compass and a current indicator are often lowered with the camera and included in the photograph. These provide indications of current velocity and direction and a means of orienting the bottom features. Locations of almost all deep water photograph stations in the Atlantic are shown on Fig. 10, and a typical deep sea camera rig is depicted in Fig. 11.

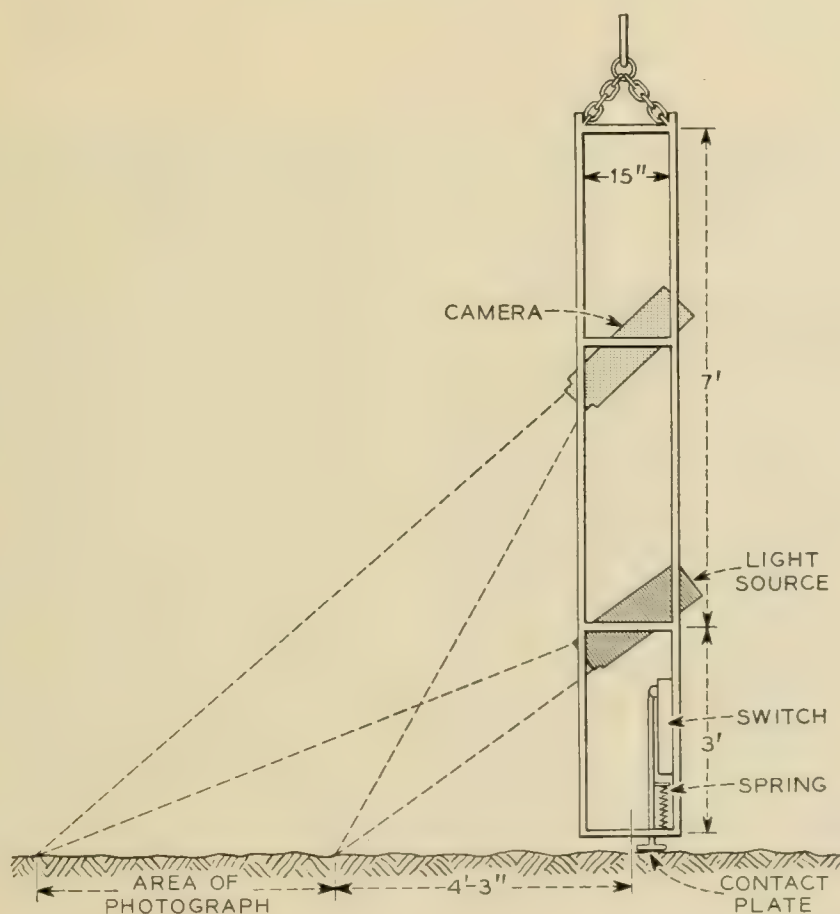


Fig. 11 — Diagram of multiple-shot underwater camera taking a bottom picture.

Television is being used in shallow waters by various organizations. However, picture resolution is poorer than that obtained photographically, and nothing has been done in depths greater than 110 fathoms.

Physical sampling is accomplished by corers, grab samplers, and rock dredgers. Earliest samples were obtained by "arming" a sounding lead with tallow, to which some of the bottom sediment would adhere. Coring is the most important source of bottom composition data. Specimens up to 70 feet in length are obtained by dropping a weighted tube vertically into the bottom sediments. About 1,200 sediment cores have been obtained in the Atlantic. The locations of most of these stations are shown in Fig. 10. Rock dredging has produced much evidence on the nature of the continental slope, the Mid-Atlantic Ridge, and on the various seamounts. Lamont's rock dredge stations are shown in Fig. 10. Not shown are a large number taken by French workers off the Bay of Biscay and off Georges Banks.

Sounding data can provide a wealth of detail in addition to the depth if an experienced operator evaluates the fathogram. The least skilled operator can differentiate between rough and smooth bottoms, while the most experienced can interpret a fathogram in terms of bottom smoothness, sediment thickness, and the location of interfaces in the sediment.

3.2 *Present Knowledge*

3.2.1 *General Characteristics*

Navigation charts sometimes include a short notation alongside a sounding, indicating the type of bottom, ranging from common terms such as sand, mud, or ooze, through the less familiar foraminifera or globigerina ooze (shells of microscopic marine life). These notations are most abundant in shallow coastal waters where they provide information for piloting and anchorages. In the less frequented depths, bottom notations on navigation charts are very rare—charts of bottom sediments commonly published are based on sparse and incomplete data and, as a result, are generalized.

Deep sea sediments have been divided into two main classes, terrigenous and pelagic. Terrigenous sediments are those derived from the erosion of the land and are found adjacent to the land masses, while the pelagic deposits are found in the deep sea and are distinguished as either organic ooze or inorganic clay. The organic oozes are composed principally of fossil remains of planktonic animals. Distribution of types of sediment is by no means static. Such factors as deposition by turbidity currents, land slides or slumps, bottom scour by ocean currents, and climatic changes continually cause changes.

3.2.2 *Sediment Densities*

Quite accurate density determination can be made by laboratory analysis of sediment from core tubes. Densities of 1.35 to 1.55 gm/cc are characteristic of the gray and red clays which cover most of the deeper parts of the ocean. The globigerina oozes which are abundant on the Mid-Atlantic ridge have densities from 1.60 to 1.75 gm/cc. Sand layers which may occur in abyssal plains at or near the sediment surface range in density from 1.65 to 2.00 gm/cc.

An increase in density as a function of depth in the core would be expected. However, after an initial increase in the first 1 or 2 meters the density usually fails to show further regular increase in cores up to 10 meters in length, and deep in the core the density often falls to within 0.1 gm/cc of the initial value.

The sediment averages about 300 fathoms in thickness over the deep ocean floor, except for the bare rock surfaces on the steeper parts of the continental slopes and submarine peaks where sediment is thin or absent. Thicknesses exceeding 100 fathoms may be reached on the continental rises.

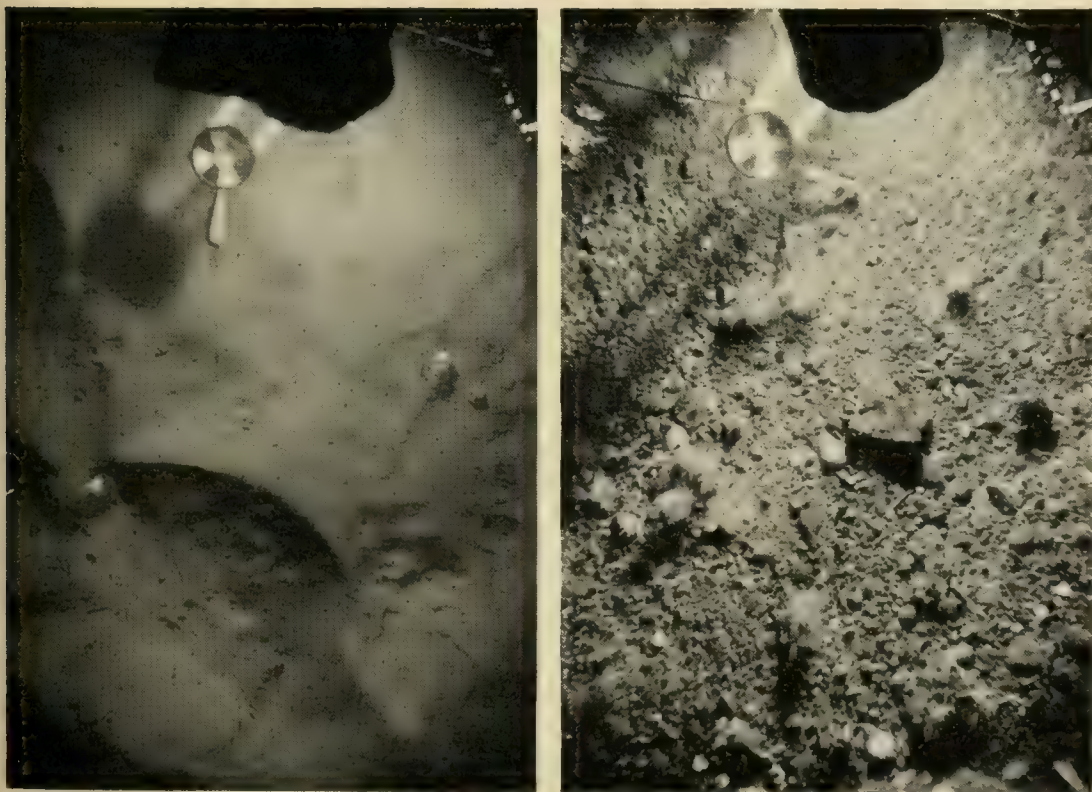


Fig. 12 — Continental-shelf photographs taken off Cape Cod. Each photograph shows an area of approximately 2.4 ft by 3.3 ft. Dials are compasses. Tassel beneath dial indicates current. Note small ripple marks in left photograph. Depths: 64 and 54 fathoms, respectively.

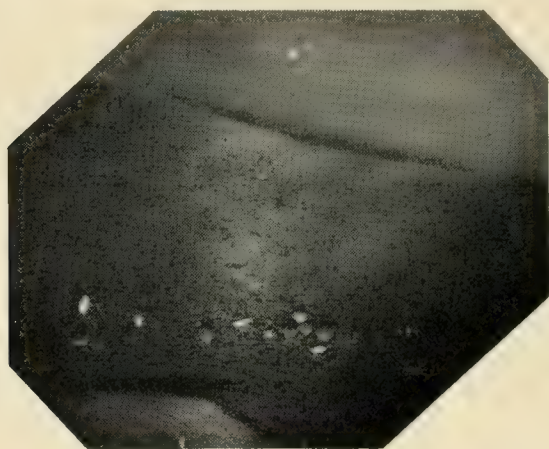


Fig. 13 — Continental-slope photograph taken at 550-fathom depth at $44^{\circ}43'N$, $54^{\circ}30'W$. Area shown in each photograph in Figs. 13–16 is about 40 square feet.

3.2.3 *The Bottom in the North Atlantic*

The continental shelf at depths less than about 70 fathoms was dry land for a considerable period prior to 11,000 years ago. Thus, the sediments of the continental shelf resembled the sediments of the coastal plain from Cape Hatteras to Cape Cod. A deposit of sand continues along the shelf edge and is generally thought to be an old beach. Landward of this is a series of irregularities generally considered to be old dunes and beaches. Photographs of the continental shelf are shown in Fig. 12.

Hardened sandstones and limestone have been recovered from the walls of submarine canyons off Georges, Browns and Banquero Banks. A rock outcrop has been photographed at a depth of 500 fathoms in a small gully south of Block Island. In other areas the continental slope is covered with low-density gray clay in which the coring rig completely buries itself. Gravel and sand form the floors of some continental slope canyons while others are deeply covered with low density mud. In many areas ancient, partly consolidated clay crops out on canyon walls.

The Western Union Company, when plowing in their continental-slope cables, had widely different experience along closely parallel cables.⁵ Presumably the differences in the depth to which the plow would penetrate were due to differences between ancient and recent compaction of sediments. Although rock was probably not encountered on these runs, it is known from dredging experience that rock can be expected.

A photograph (Fig. 13) of the bottom at 550 fathoms depth south of the Grand Banks reveals huge ripple marks. It is not difficult to imagine that cable chafe would be appreciable in such an area.

Beneath the nearly flat abyssal plains alternate layers of sand, silt,

and red clay make up the approximately 300 fathoms of sediment. A photograph (Fig. 14) shows the bottom at a depth of 3,000 fathoms in the abyssal hills. This remarkable shot shows ball-shaped objects that have been identified as manganese nodules. The most interesting feature of this photo is the scour marks around the objects, implying an appreciable current at a depth of 3,000 fathoms.

Seamounts present extremely varied conditions. Rocks, from crystalline basalt through hardened limestone to soft marl, are encountered. Sediments, including sticky ancient formations, deep sea oozes, and shell sand are found. Photographs show all types from tranquil mud bottom through wave-rippled silt and sand to craggy rock. Some of these types are illustrated in Fig. 15.

The flanks of the Mid-Atlantic Ridge are areas of irregular topography. The steeper slopes are probably bare rock and the sediment removed from these slopes probably is deposited in the intermountain basins. Cores are usually of globigerina ooze, but the different rates of sedimentation on steep slopes and on basin floors cause changes in thickness of sediment and relatively great changes in physical properties over short distances. The deeper flanks of the ridge are covered by red clay. A very similar bottom is found on the Bermuda Rise.



Fig. 14 — Abyssal-hills photograph taken at 3,190-fathom depth at 29°17'N, 57°22'W.

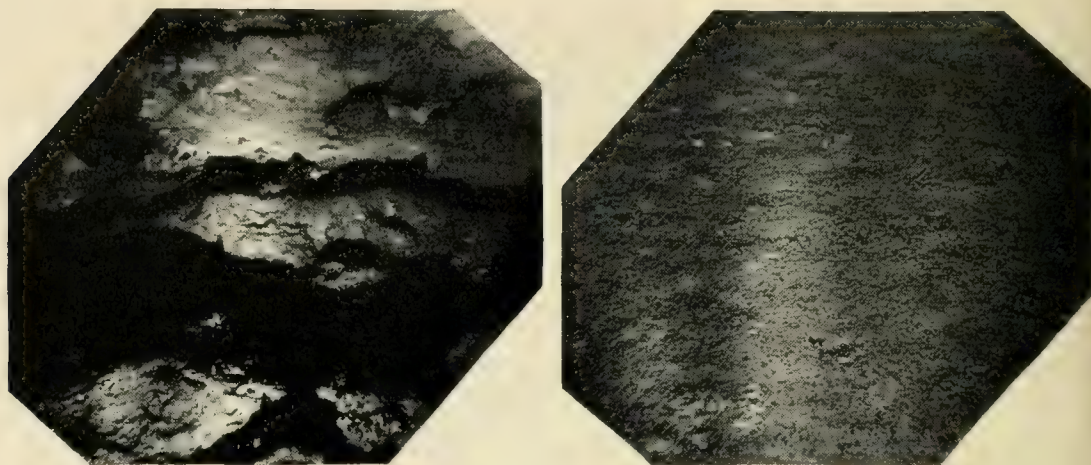


Fig. 15 — Seamount photographs taken near summit of seamount on Bermuda Rise at $34^{\circ}38'N$, $56^{\circ}53'W$ at depth of 1,370 fathoms. The two photos were taken about 100 feet apart, indicating the rapid alternation of ooze and rock bottom over short distances.

The crest of the Mid-Atlantic Ridge is similar, as a bottom type, to the seamounts previously described. Dredge hauls have brought up mostly basalt, although a few fragments of limestone have also been retrieved. The photographs shown in Fig. 16 were taken about 60 feet apart on the Mid-Atlantic Ridge. They illustrate a change from smooth to rocky conditions in this short distance.

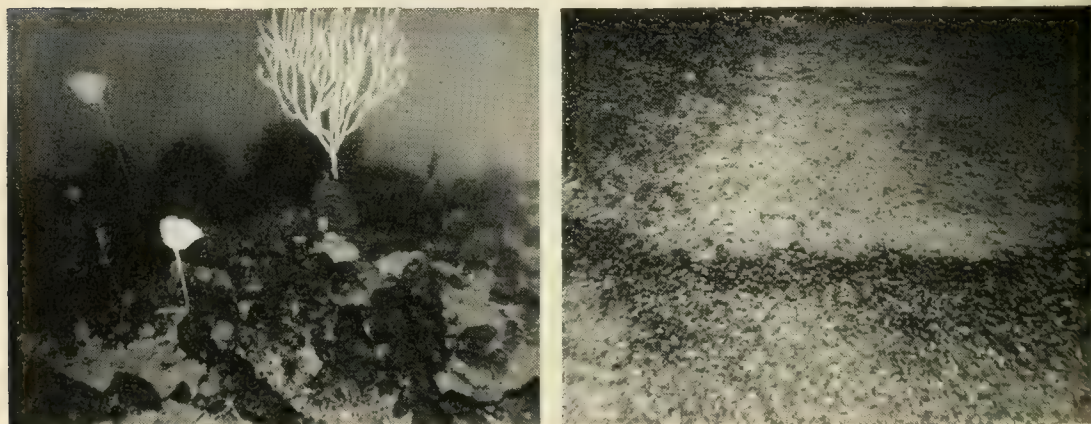


Fig. 16 — Mid-Atlantic Ridge photographs (1,500-fathom depth at $48^{\circ}30'N$, $28^{\circ}48'W$). The two pictures were taken about 60 feet apart. Out of 60 photographs taken at similar intervals in this location three were similar to that on the left and the remainder resembled that on the right. The dark band in the right-hand picture is probably composed of gravel and sand of the dark-colored rock of the peaks, while the white underlying layer is clay or ooze. The dark band was produced by a current which swept the dark material over the light-colored material.

IV. TEMPERATURES*

4.1 *General*

The temperature of a given point on the deep sea floor is determined by the system of ocean circulation. Study of deep sea circulation is still in an early stage and theories which would permit prediction of changes are still in a rudimentary form. In addition, observations of actual bottom temperature are few. Thus, a study of the temperature environment of submarine cables must proceed by evaluating those data that exist and striving for increased understanding of the underlying circulation processes.

The new electronic thermometer under development at the Lamont Geological Observatory determines temperature with an accuracy of 0.01°C by the frequency of an oscillator employing thermistors in an RC network. The oscillator is lowered on the end of a cable and its frequency is monitored by equipment installed aboard ship.

Bottom-temperature changes might be predicted if the rate and direction of circulation of the sea water could be determined. This is being studied by sonar tracking of a submerged blimp designed so that it has negative buoyancy at the surface but is neutrally buoyant at the level where the measurement is desired. This method has been used down to depths of 3,000 meters. Results have indicated much higher velocities than hitherto suspected. Near the base of the continental slope off the eastern United States near-bottom currents of $\frac{1}{3}$ knot have recently been observed by this method.

Another method depends on the measurement of the time elapsed since a given water mass was at the surface by radiocarbon dating of sea-water samples. At the surface, water is in free exchange with the atmosphere and acquires a radiocarbon concentration in equilibrium with that of the atmosphere. As the water sinks from the surface to enter the deep sea circulation system, it is cut off from the supply of fresh radiocarbon, and radioactive decay reduces the content of Carbon 14 at a rate given by its half life. Thus, the measurement of the radiocarbon content of a given sea water sample ideally will give the time at which this sample left the surface of the ocean.

* In this section dealing with temperature, depths are given in meters rather than in fathoms. Temperature has largely been of interest in physical oceanography where volume is of concern. This has led to the use of meters. Since a nautical mile is approximately 1,000 fathoms, the fathom has been widely used in topographic work. One fathom equals approximately two meters.

4.2 *Characteristics of Available Information*

4.2.1 *Sources of Data*

Observations of temperature in the deep sea were first made in the mid-nineteenth century. The early observations were made with crude instruments and are now of purely historical value. The *Challenger* expedition of 1872–1876 made several hundred observations but with maximum-minimum thermometers unprotected from pressure. In the late nineteenth century the Richter reversing thermometer was invented and by the turn of the century they were used by nearly all scientific expeditions. Cable ships have taken many observations but almost always with maximum-minimum thermometers.

Major expeditions have published volumes which included tabulated lists of temperature, salinity, oxygen, etc., for each station occupied, while the shorter expeditions and those institutions which continually collect oceanographic data in the North Atlantic publish their observations in the “Bulletin Hydrographique,” a journal published by the International Commission for the Exploration of the Sea, Copenhagen. In addition, unpublished data are available from the files of oceanographic institutions.

Expedition reports give estimates of the reliability and accuracy of their data and usually describe the calibration tests used to determine the accuracy. The “Bulletin Hydrographique” merely publishes the data without comment. The scarcity of data and the tendency to systematic errors in single sets of data coupled with the temperature changes now being demonstrated for deep ocean water masses tend to frustrate efforts to evaluate the accuracy of data.

4.2.2 *Methods of Evaluation*

Evaluation of temperature data involves comparison of nearby observations, verification of the original data sheet, and checking for errors in computation. The calibration of the thermometers is ordinarily done with great care, and observations are generally accurate to $\pm 0.05^{\circ}\text{C}$. When the thermometer fails to function properly, the temperature is usually so far off that the observation is not reported. The main error comes in the determination of depth of observations. The length of wire paid out to reach a stated depth varies with the magnitude of winds and currents in the area. On early expeditions this led to large errors in observation. The use of two thermometers, one in a pressure case and one unprotected from pressure, allows the calculation of depth of observations with relatively great accuracy.

One method of understanding temperature changes is to consider the movement of the water masses. These movements depend on density gradients which are directly dependent on the salinity and temperature. This type of study is applicable in shallow water where circulation of the surface layers is brisk. For the deep sea, however, dynamic calculations are ambiguous; different investigators using the same data not only arrive at different values for velocity but often arrive at opposite directions of flow. Thus, continuity considerations and the conservation of volume are the primary factors in studies of the deep-water circulation. It is important when evaluating temperature data and studying temperature change and rate of bottom water circulation to study the entire process and not be limited to temperature observations, even though the temperature is the desired final answer.

4.3 *Temperature in the North Atlantic*

4.3.1 *General*

Ocean-bottom temperatures in shallow water (less than 200 meters deep) are affected by seasonal air temperatures and movements of local masses of water. Thus, each specific area exhibits its own pattern of temperature changes. Data are fairly abundant in shallow water areas, so that despite the large and often erratic changes, it is generally possible to determine roughly the bottom-temperature cycle from existing data. However, the local nature of the phenomena makes it desirable to concentrate any detailed studies on areas of immediate interest rather than to attempt broad generalizations.

In deep water (depth greater than 200 meters) bottom temperatures and their variations result from large-scale, oceanwide topographical and circulation phenomena. At the same time, the paucity of data makes an analysis of any particular locale quite difficult. Thus, a general study of deep-water bottom temperatures in the North Atlantic presents the best hope of obtaining at least some useful data.

4.3.2 *Shallow Water (Depth less than 200 Meters)*

In many shallow-water areas near shore (depth less than 50 meters) there are predictable seasonal changes in temperature of the order of 10°C. In harbors and bays where interchange with open ocean water is restricted the seasonal temperature curve will approximate the seasonal air-temperature curves except that the amplitude of the sea bottom temperature changes will be smaller and the peaks and troughs will be

slightly retarded. In the open water of the continental shelf there is often a strong stratification of water masses and the vernal cycle is either strongly retarded (by several months) or completely obscured.

On the open shelf the bottom-temperature cycle is controlled by shifting currents and wedges of water which flow in along the bottom from the open ocean. In some areas these changes go through essentially the same cycle each year. The Irish Sea and the continental shelf west of Scotland are areas where the cycle is so regular that one can safely predict the bottom temperature for a given month within an accuracy of ± 1.5 – 2°C . On the other hand, the bottom temperature on the Grand Banks changes radically from day to day and week to week. It is possible to show that on the Grand Banks summer temperatures are on an average 2° colder than winter bottom temperatures. The day-to-day temperature changes can amount to 5°C or more.

In one particularly well-studied area off Halifax, Nova Scotia, an interesting complication has been discovered. In this 10,000-square mile area the bottom temperatures had been studied for about twenty-five years, observations having been taken at different times of the year. The bottom temperature was considered known to 1°C . It has more recently been found that on occasion, the 8°C water is displaced upward by an underflow (incursion) of 2°C water which suddenly lowers the bottom temperature by 6°C . However, after a few weeks the bottom temperature again approaches the usual value of 8°C . Such incursions of contrasting water (both cold and warm) from the open ocean are as yet only partially understood.

The maximum amplitude of temperature changes in bays and near shore areas of both seasonal and erratic nature often approaches 20°C . On the outer shelves 8°C would be the maximum change expected.

It is now well established that the average temperature over wide areas in the North Atlantic has undergone a gradual increase for the past one hundred and fifty years. The average bottom temperature on the Nova Scotian shelf increased 2°C between the early and mid-nineteen thirties and the late nineteen forties. Similar changes have been reported for the area near Iceland. At present, no sure way of predicting the future longterm changes of temperature is available. Changes of 1° per decade may be experienced.

4.3.3 *Deep Water (Depth more than 200 Meters)*

A search was made for all deep sea temperature measurements taken with accurate thermometers in depths greater than 2,000 meters, from



Fig. 17 — Accurate deep-water temperature measurements in depths greater than 2,000 meters. Points marked by solid circles indicate observations within 100 meters of the bottom.

the Equator to the Iceland-Faeroe-Greenland Ridge (which divides the Atlantic Ocean from the Norwegian Sea). Approximately 600 observations (Fig. 17) were found after a search of all data up to 1950 and much of the data up to 1954. (The "Bulletin Hydrographique" is 5 to 6 years behind in publication.) Since the temperature at intermediate depths is of interest to the cable engineer only to the extent that he can use it to determine bottom temperatures, the observations in depths greater than 2,000 meters which lay within 100 meters of the bottom were sorted out. Only about 150 observations (Fig. 17) were found in this

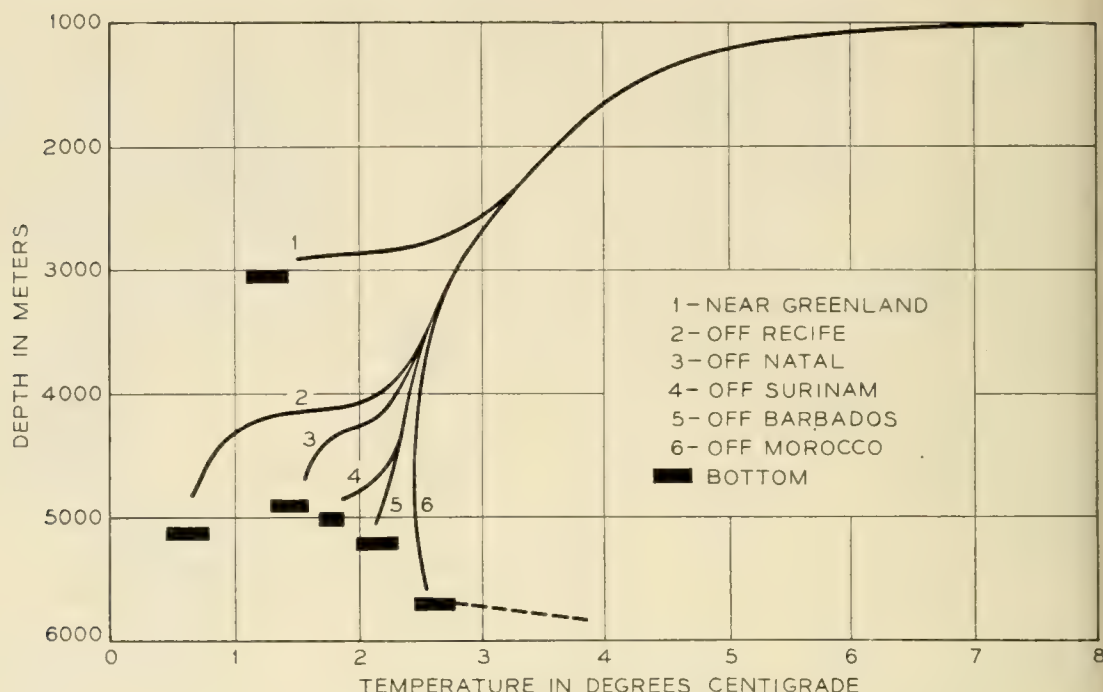


Fig. 18 — Deep-sea temperature gradients. Mean gradient in the sediments is indicated by the dashed line.

category and most of these were either between Greenland and Labrador or near the Equator. The number of actual bottom observations is limited by inability to determine when the bottom was reached and a reluctance on the part of observers to risk losing expensive equipment by having it snagged on the bottom.

At the present time there are insufficient data or knowledge of the mechanism involved to permit reliable extrapolation of bottom temperatures from a series of mid-depth observations. Fig. 18 illustrates the problem. Here are six different near-bottom gradients observed in different parts of the Atlantic. Assuming these gradients terminated 500 meters above the bottom (as do many of the observed data), it is apparent that extrapolating such data to the bottom is not feasible. Measurements of gradients to the bottom at stations for which near-bottom data exist, coupled with a knowledge of the processes causing the gradient, may make it possible in the future to make use of many of the old mid-depth observations in studying bottom temperature.

From the compilation of available data, the three profiles (Figs. 20-22) whose locations are shown in Fig. 19 were prepared. More recent studies indicating that deep water temperatures may vary a few tenths of a degree Centigrade with time make it probable that some of the ripples in the isotherms are not real but instead reflect the fact that the data were taken at widely different times. The data for Profile 1 were

taken in the same year and do not show these effects. The large fluctuation on Profile 3 is probably due to variations with time.

At a certain depth, often about 1,200 meters, a sharp change in temperature takes place. An interface known as the main thermocline separates the warm surface waters from the cold (2° – 4°C) deep water. This boundary shifts slightly from time to time and is affected by subsurface (internal) tidal waves. Thus, cable laid near the main thermocline may undergo temperature changes of a few degrees. These changes may be of different periods depending on their causes, e.g., subsurface waves, seasonal changes, or long term changes in ocean regime. It is not known if a long-term change in depth of the thermocline has occurred but such change seems probable.

The greatest percentage of any transatlantic cable route lies in depths greater than 2,000 meters and thus the temperature changes in the deep sea are of primary importance to the engineering of a cable. The temperatures are low, averaging 3°C , and the temperature changes are

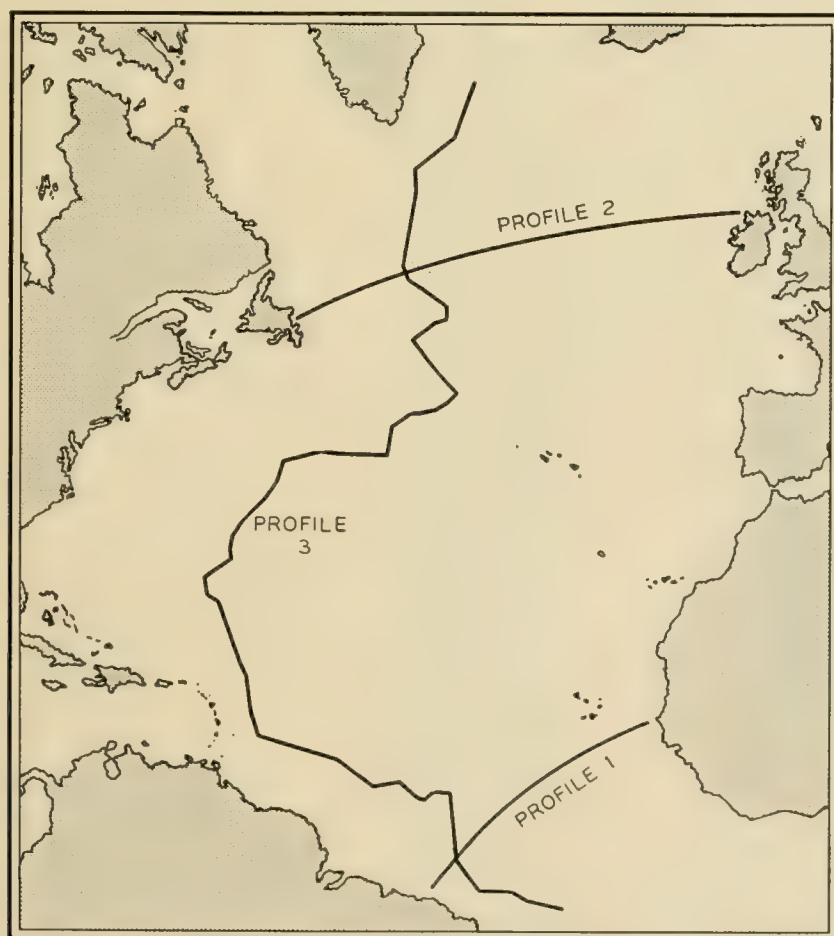


Fig. 19 — Positions of Temperature Profiles 1, 2, and 3.

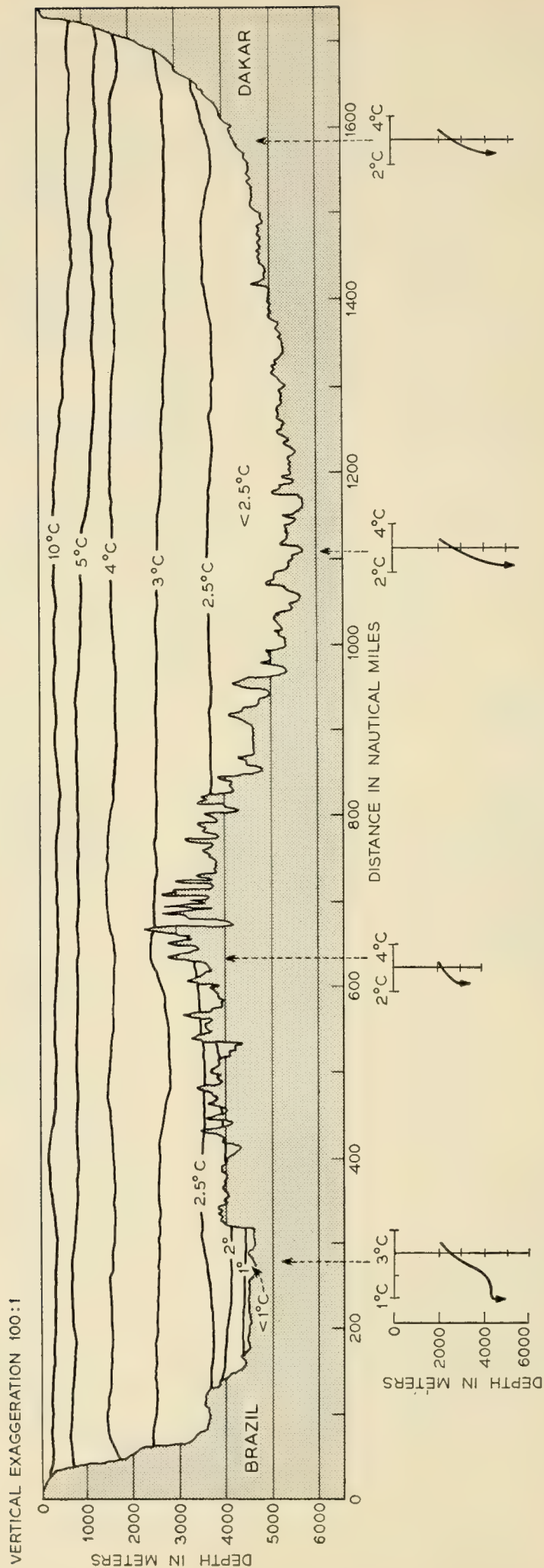


Fig. 20 — Profile 1, Northeast Brazil to Dakar, Africa. Isotherms in degrees Centigrade. Data from 1928.

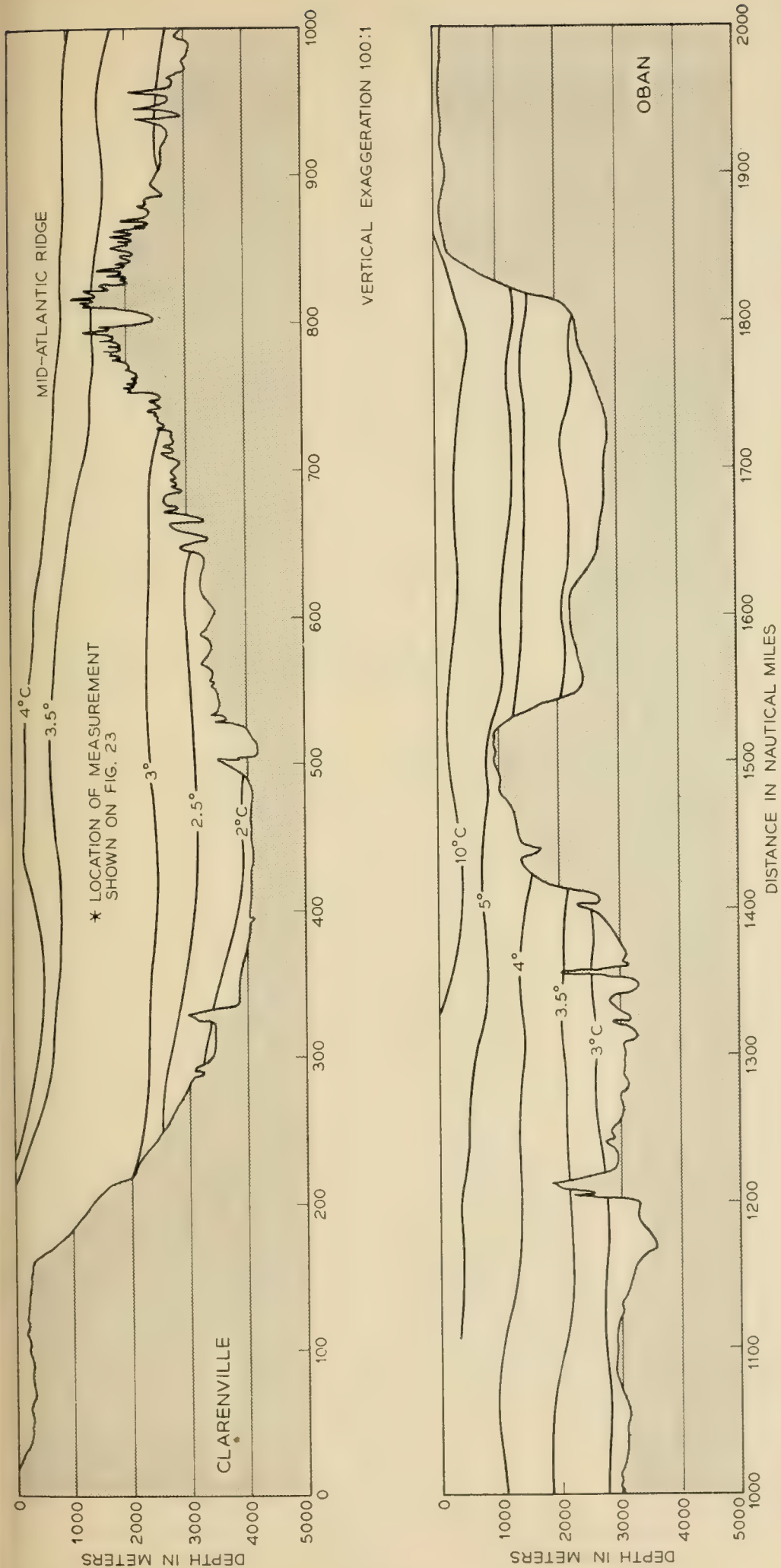


Fig. 21 — Profile 2, Clarenville, Newfoundland, to Oban, Scotland. Isotherms in degrees Centigrade. Data from several years.

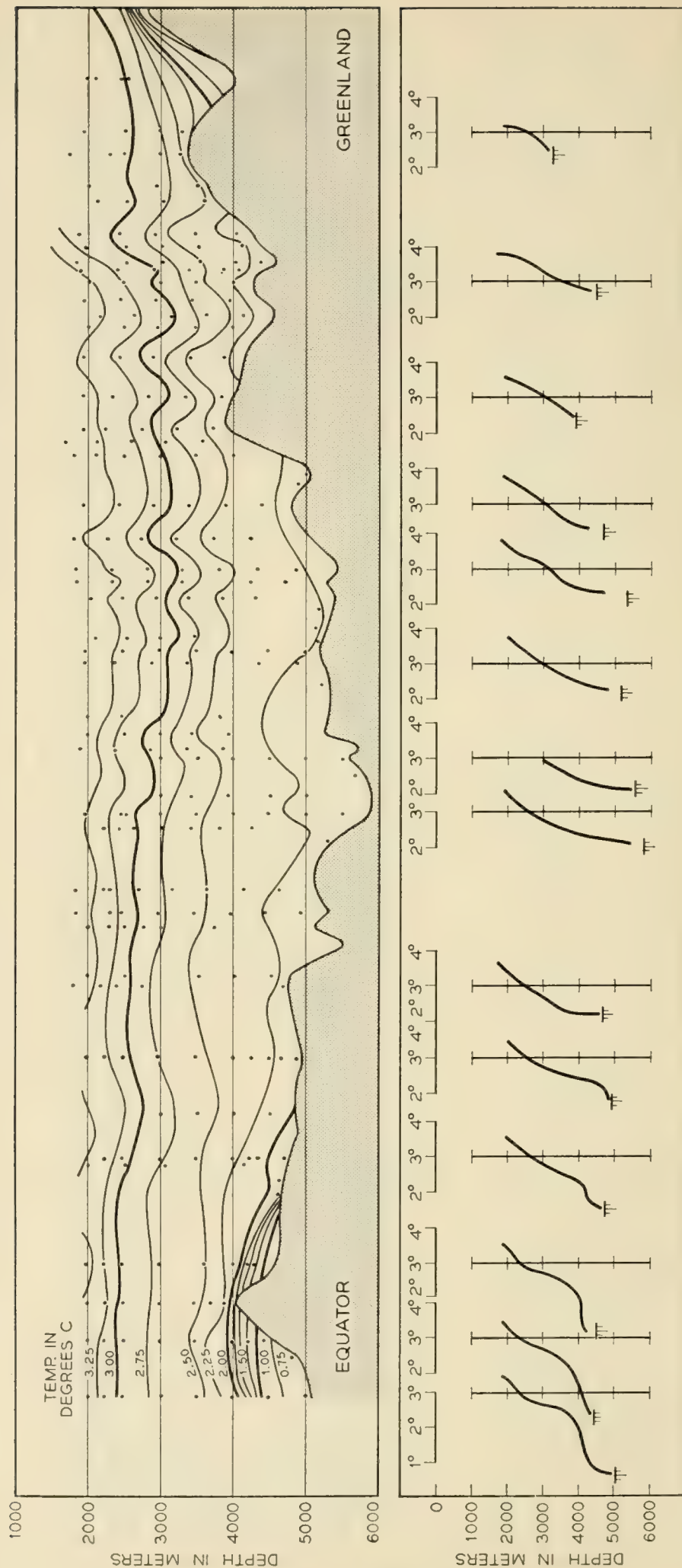


Fig. 22 — Profile 3, Equator to Greenland west of Mid-Atlantic Ridge. Top, isotherms in degrees Centigrade; bottom, gradients at corresponding locations. Data from 1920 to 1955.

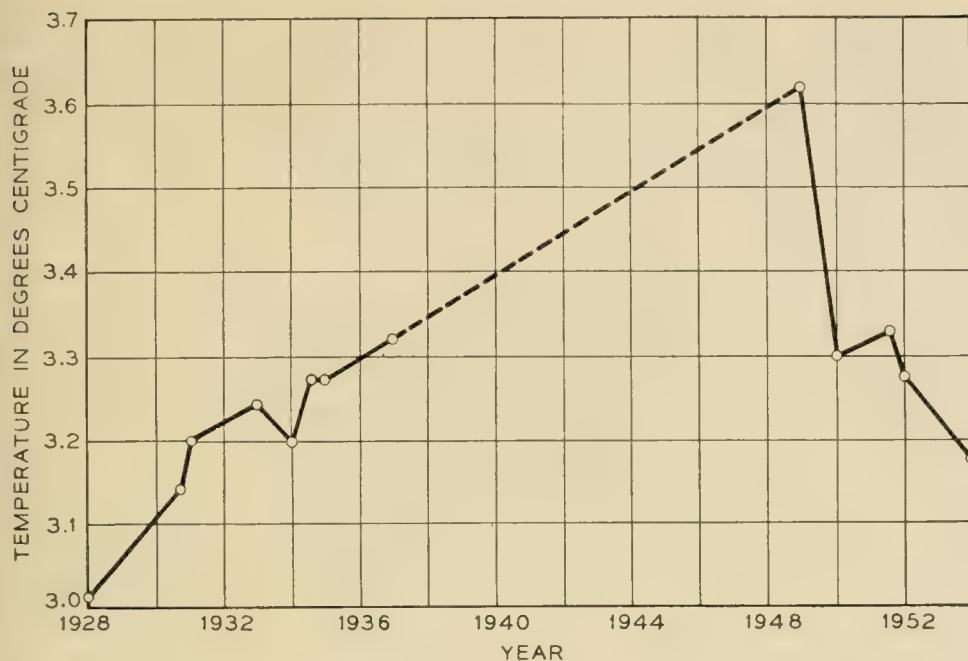


Fig. 23 — Variation of mean temperature at 1,500-meter depth in the southern Labrador Sea at about the point indicated in Profile 2, Fig. 21.

probably small. At one time it was assumed that the ocean temperature in depths exceeding 1,000 meters remained constant. As more information is gathered it is becoming evident that changes in temperatures have occurred in deep water, but neither the mechanism nor the time scale of the changes is as yet completely understood.

Due to the virtual absence of deep sea bottom temperature observations it is not possible to determine changes in temperatures by comparison of repeated measurements at approximately the same position. In a few limited areas repeated observations have been taken to depths of about 3,000 meters. All observations for one such area northeast of Newfoundland have been studied in search of long-term trends. It has been found that the water temperature between 500 and 2,000 meters in this area is nearly constant, both vertically and laterally over a wide area during any one year. It is thus meaningful to compare the temperature at 1,500 meters depth for a series of years. Fig. 23 shows the results of this comparison which indicates a 0.6°C increase between 1928 and 1949, and a 0.4°C drop from 1949 to 1954. This curve resembles the air- and sea-surface temperature averages for Atlantic-coast stations during the same years.

In the area between Labrador and Greenland a moderate number of near-bottom temperature measurements have been made since 1928. No systematic curve can be drawn but it seems fairly certain that bot-

tom-temperature changes of 1°C have occurred in this area. These temperature changes are apparently caused by cold water cascading down the continental slope from the shelf off Greenland. It can be presumed that this water will flow south along the deep ocean basin. Depending on its velocity it will be more or less displaced towards the western margin of the basin. As the water masses flow south, lateral mixing should reduce the temperature contrast with the surrounding water.

In one area south of Newfoundland a study of scattered temperature observations made since the year 1900 indicates that water above 3,000 meters has shown a temperature increase and that the water below 3,000 meters has gradually decreased in temperature. The decrease at the bottom at a depth of 5,000 meters appears to have amounted to 0.2°C in fifty years and the maximum increase at depths less than 3,000 meters to about 0.5°C . In the high latitudes of the North Atlantic, temperature changes probably rarely exceed 1°C at depths greater than 2,000 meters and changes of a few tenths of a degree are more common.

V. CATASTROPHIC CHANGES IN THE OCEAN BOTTOM

5.1 *Earthquakes*

Earthquakes may cause damage to submarine cables by triggering the movements of rock and sediments (turbidity currents) and possibly through the effect of the actual earth vibration itself. The most serious threat to cables arises not from the direct effect of the earthquake's energy on the cable but from its ability to generate slumps, slides, and turbidity currents. It can be shown that, in an area of high seismic activity, the slopes will be nearly bare of loose sediment so that earthquakes in such an area will not result in gravitational displacements of sufficient size to cause serious damage to cable.

However, in areas such as the continental slopes of the Atlantic where few shocks have been recorded and where loose sediment has therefore accumulated, it can be expected that any quake of moderate or even small size will be sufficient to generate a turbidity current. Thus, quakes on continental slopes adjacent to cable routes are likely to be extremely destructive. There is no known method of predicting where earthquakes will occur *outside* the major seismic belts.

Fig. 24 shows the distribution of earthquakes in the North Atlantic. The North Atlantic is the best-monitored ocean because of the extensive network of seismograph stations closely adjacent in North America and Europe. All earthquakes reported for the Atlantic Ocean between 1910 and 1956 have been compiled and plotted together with the best bathy-

metric information. The study resulted in conclusions which are of great geological significance, as well as of interest in cable engineering. It was found that the narrow earthquake belt which runs the length of the ocean accurately follows a median trench or rift in the central highland zone of the Mid-Atlantic Ridge, see Fig. 25. (The Mid-Atlantic Rift

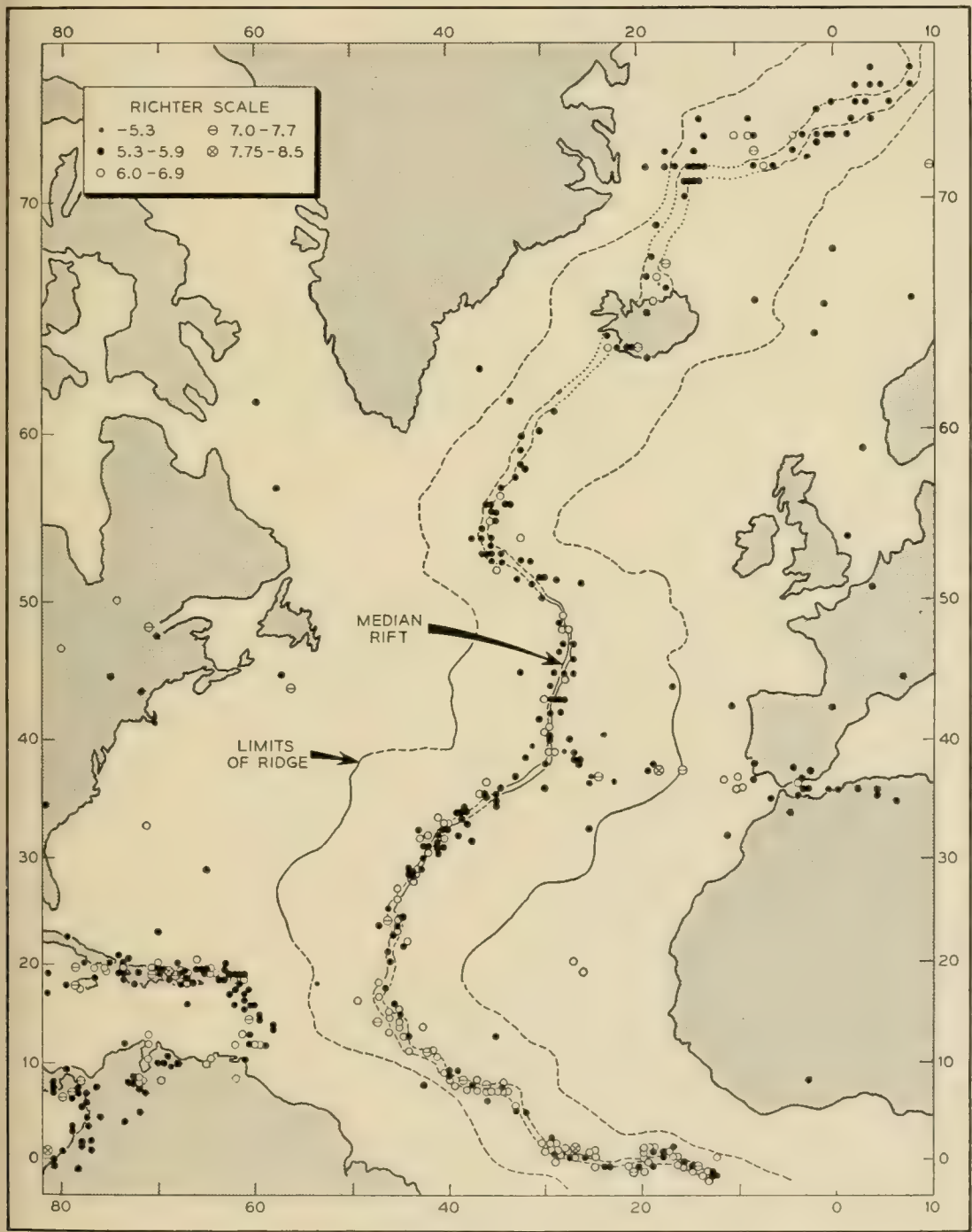


Fig. 24 — Earthquake epicenters in the North Atlantic 1910-1956. Magnitudes according to the Richter scale.

seismic belt is of secondary importance when compared with the intense seismic belt associated with the deep trenches of the Pacific.)

There is a second belt which extends from the Azores to the Iberian Peninsula, and there is a major earthquake belt which follows the West Indies Island Arc along part of the western boundary of the Atlantic. In addition to these belts there are a few scattered earthquakes around the margin of the ocean basin.

The locations of earthquakes in the ocean are in general poorly determined, and an accuracy of $\pm \frac{1}{2}^\circ$ (~ 30 miles) is about all that is ever claimed. Within this measurement accuracy, all quakes in the central

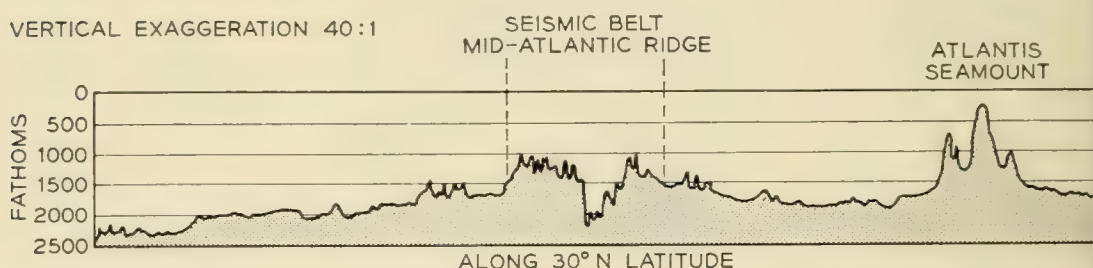


Fig. 25 — Profile of Mid-Atlantic Median Rift showing maximum limits of seismic belt.

part of the Atlantic fall in the Mid-Atlantic rift zone. Thus, there is a strong probability that all quakes on the Mid-Atlantic Ridge fall in this median trench. (Earthquakes of small magnitude are often detected when the signals are not sufficiently clear to allow a position to be determined. There are undoubtedly other quakes that are altogether undetected.) All cables to Europe must cross this rift zone and thus will undoubtedly undergo earthquake shocks from time to time. There is no positive evidence that quakes have broken telegraph cables where they cross the rift. There are, however, several cases where the telegraph cables have parted in or near the rift.

In the western Atlantic there have been recorded only six earthquakes that fall outside the belts described above. Two of these quakes occurred on the Bermuda Rise well away from cable routes. One quake occurred on the continental slope south of Newfoundland (the famous destructive Grand Banks earthquake); a second small shock nearby apparently caused no damage to cables. Two more quakes were centered off the Labrador Coast. Many cases have been recorded along the Pacific coast of Central and South America where the cables failed in deep water following an earthquake.

5.2 *Turbidity Currents — Slumps*

The turbidity current is a flow of sediment-laden water which occurs when an unstable mass of sediment at the top of a relatively steep slope is jarred loose and slides down the slope. As the slide or slump travels down the slope, more and more water becomes mixed in the mass, giving it high fluidity combined with its high density. The currents reach very high velocities, enabling them to spread for vast distances across the abyssal plains. Such slides occur at the edge of continental slopes, in particular, in the vicinity of river mouths, off prominent capes, or near banks. They are triggered by earthquakes, hurricanes, or floods.

Three excellent pieces of evidence support the turbidity-current theory and all have important implications for submarine cables. The turbidity-current hypothesis was first introduced to explain the erosion of submarine canyons in the continental slopes. Supporting evidence was found in submarine cable breakage following the Grand Banks earthquake of 1929⁶ and the Orleansville (Algeria) earthquake of 1954.⁷ In both cases, submarine cables were broken consecutively in the order of increasing distance *downslope* from the epicenter of the earthquake. All of the cables were broken in at least two widely separated places (100 miles or more apart). The sections between breaks were swept away and/or buried beneath sediment deposited by the current so that repair ships were never able to locate a large proportion of these sections. Fig. 26 shows the area of the Grand Banks turbidity current.

After study of this evidence, Heezen and Ewing concluded that the area covered by the current would show graded sediments as a result of deposition by the turbidity current. This was substantiated by the evidence of cores obtained from the locations shown on Fig. 26.

Cable damage resulting from turbidity currents due to slumps at a river mouth is well illustrated by that off the mouth of the Magdalena River (Columbia, S. A.). On August 30, 1935, the disappearance of 480 meters of the western breakwater and most of the river bar resulted from a slide which produced a 36-foot channel across the bar. The same night, tension breaks occurred in the submarine cable from Barranquilla to Maracaibo, which crosses the submarine canyon 15 miles off the mouth of the Magdalena in a depth of about 700 fathoms. The cable, when brought up for repair, was tightly wrapped with green grass of a type which grows in the marshes near the jetties.

The same pattern has been repeated 15 times since the cable was laid in 1930. The breakage of the cables has occurred most frequently in August and late November to early December; the two periods of the

highest water of the river and of the strongest trade winds — conditions favorable for the triggering of turbidity currents. A chart of the area showing the cable breakage on the steep continental slope is shown in Fig. 27.

Fig. 28 shows areas of the world inaccessible to shallow-water-derived turbidity currents, because of their elevation or because of protective

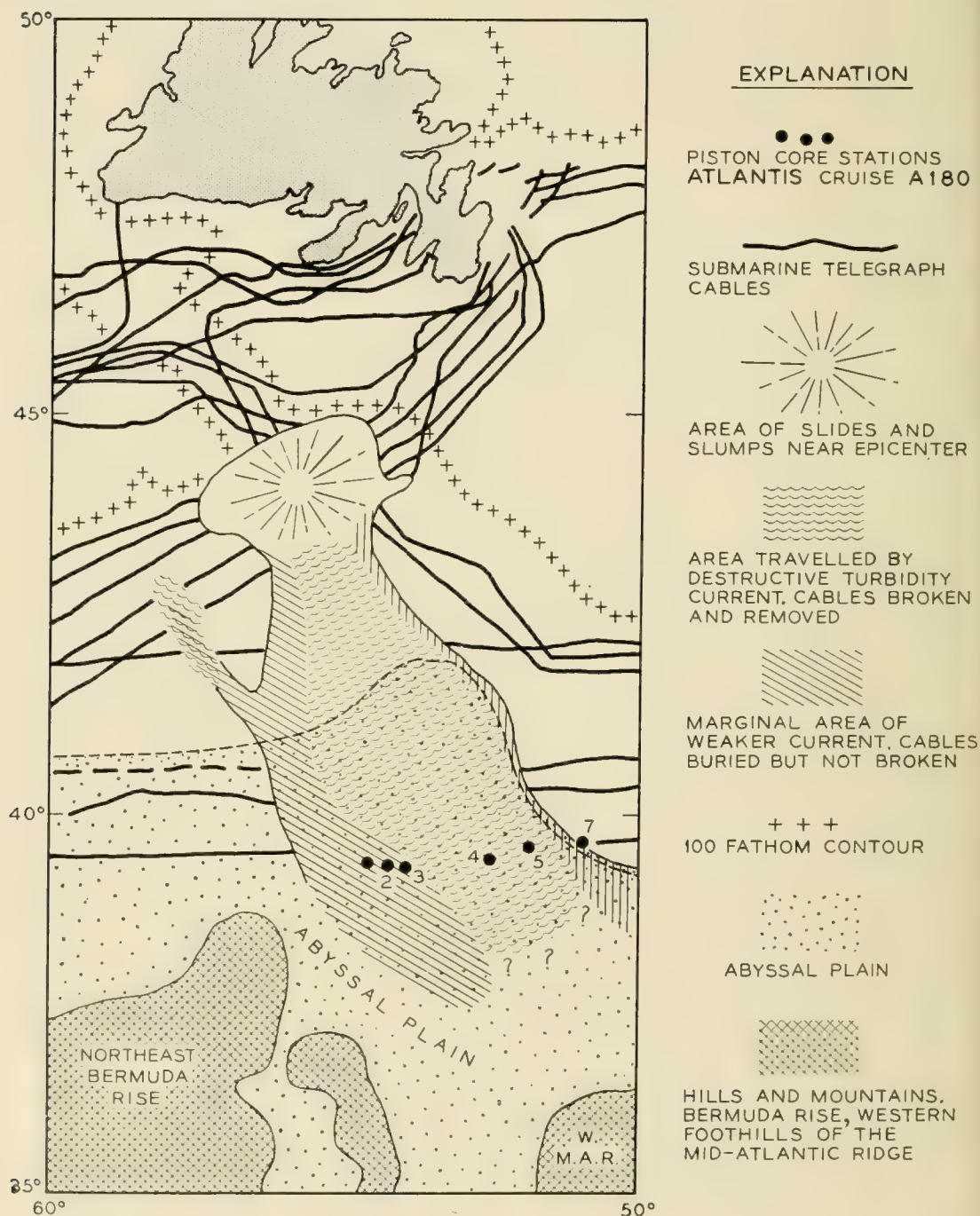


Fig. 26 — The 1929 Grand Banks turbidity current.

barriers. The arrows show some modern turbidity currents which have been documented through the breakage of telegraph cables. This chart is not yet complete and it is expected that many more arrows can be added, each representing one or more currents, and some representing as many as 20 documented cases.

Areas which have experienced turbidity current flows in the past can be determined by a careful study of the nature of the bottom. From a study of sediments and topography of areas along a proposed route it is possible to delimit areas where turbidity current could and could not

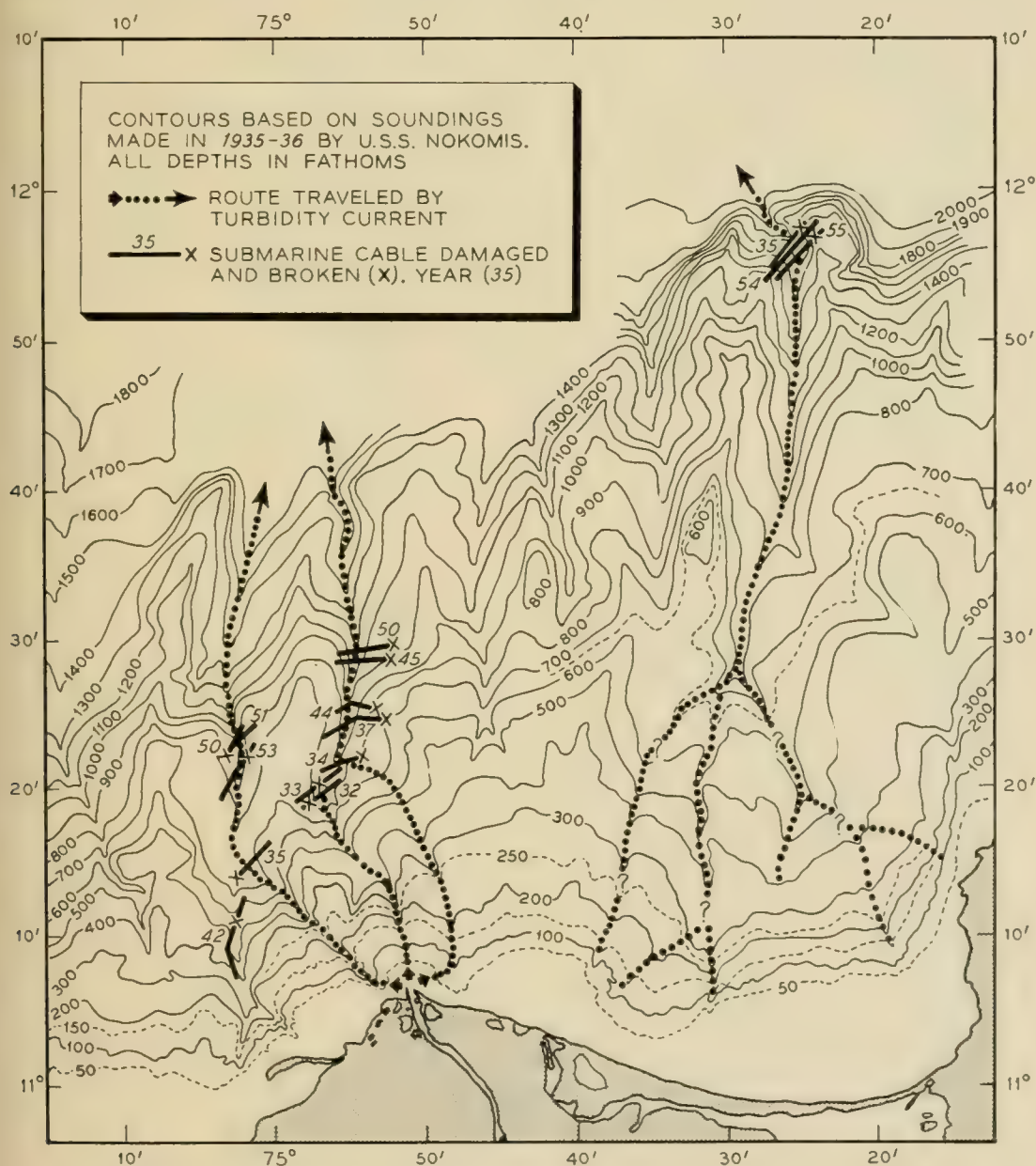


Fig. 27 — Cable breaks in the submarine canyons off the Magdalena River, Colombia, South America.

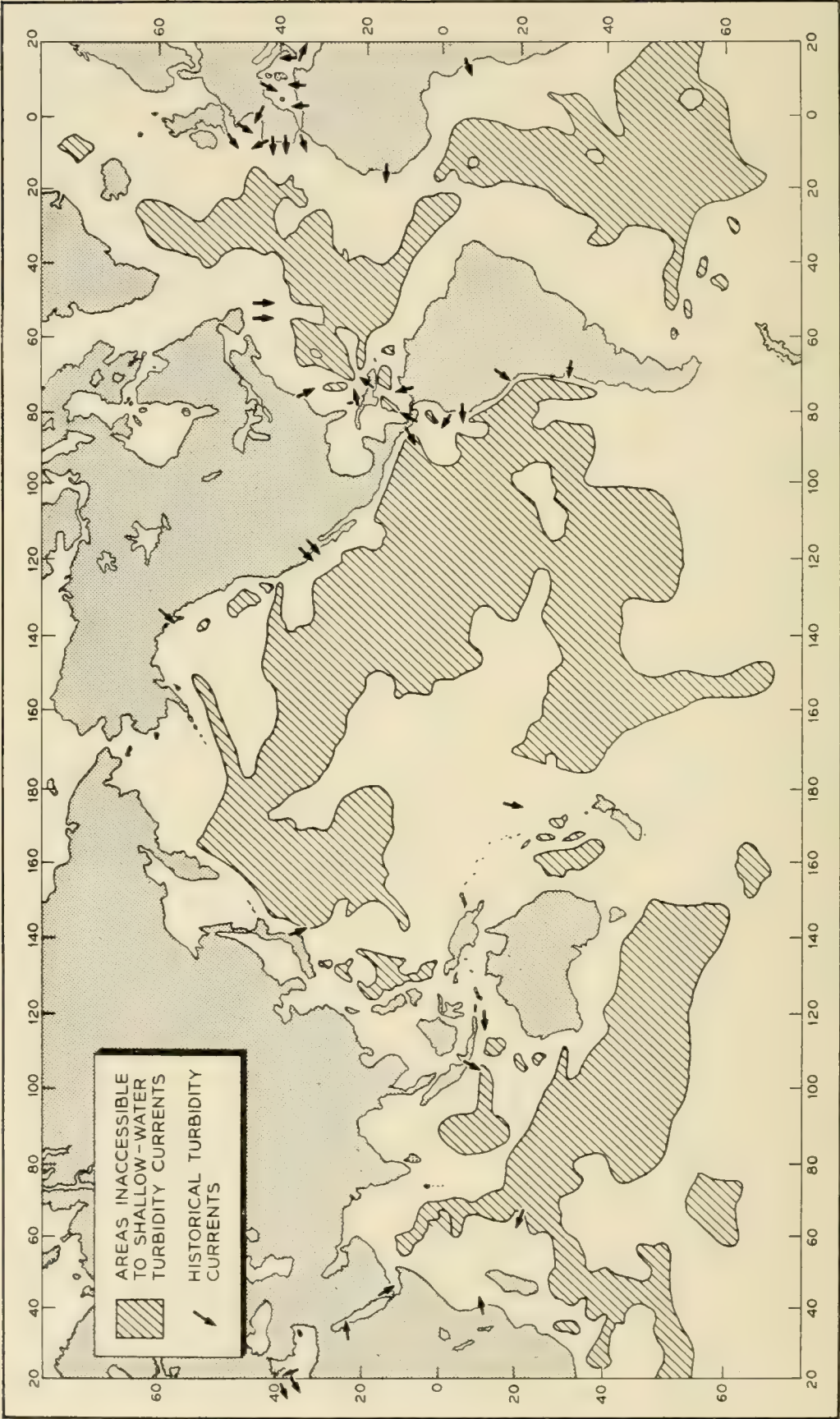


Fig. 28 — Areas of the world inaccessible to turbidity currents from shallow water.

occur, and in some cases to outline areas where they occur at frequent intervals. If, after the completion of such a study, it is decided to take the risk of laying a cable across a dangerous area and at a later date a turbidity current does occur, it will be possible to predict the length of cable to be replaced under various conditions.

VI. CONCLUSIONS

The application of the rapidly developing science of oceanography to the development and engineering of submarine cable systems will permit refinements in design, cable placing techniques, route selection, and repair operations. Existing knowledge, when evaluated and codified, provides many useful data for immediate application. Continuing study of topography, the nature of the bottom, causes of cable failures, and deep sea circulation will permit further advances in the engineering of submarine cable systems. It is to be expected that the value and usefulness of submarine cable oceanographic studies will be substantially extended as geologic and oceanographic researchers broaden the understanding of the natural laws and processes which govern and produce the ocean-bottom environment. Through such knowledge, the data of many fields can be coordinated, permitting better explanation of past events and more accurate prediction of future conditions.

ACKNOWLEDGEMENTS

The authors wish to acknowledge the contributions of Marie Tharp, W. Fedukowicz, and H. Foster of Lamont Geological Observatory, and A. L. Hale and H. W. Anson of Bell Telephone Laboratories. The encouragement and guidance of Dr. Maurice Ewing has been of great value.

REFERENCES

1. E. E. Zajac, Dynamics and Kinematics of Laying and Recovery of Submarine Cable, pp. 1129-1208, this issue.
2. L. R. Snoke, Resistance of Organic Materials and Cable Structure to Marine Biological Attack, pp. 1095-1128, this issue.
3. B. Luskin, B. C. Heezen, M. Ewing, and M. Landisman, Precision Measurement of Ocean Depth, *Deep-Sea Research*, **1**, pp. 131-140, April, 1954.
4. D. J. Matthews, *Tables of the Velocity of Sound in Pure Water and Sea Water for Use in Echo Sounding and Sound Ranging*, Admiralty, London, 1939.
5. C. S. Lawton, The Submarine Cable Plow, *A.I.E.E. Trans.*, **58**, p. 685, 1939.
6. B. C. Heezen and M. Ewing, Turbidity Currents and Submarine Slumps, and the 1929 Grand Banks Earthquake, *American Journal of Science*, **250**, pp. 849-873, Dec., 1952.
7. B. C. Heezen and M. Ewing, Orleansville Earthquake and Turbidity Currents, *Bull. Am. Assoc. Petroleum Geologists*, **39**, pp. 2505-2514, Dec., 1955.

Resistance of Organic Materials and Cable Structures to Marine Biological Attack

By LLOYD R. SNOKE

(Manuscript received June 6, 1957)

The increasing use of submarine telephone cable has resulted in the need for information on the performance of organic materials and cable structures under marine conditions. Recently, Bell Telephone Laboratories initiated a program to acquire fundamental data on the resistance of a wide range of organic materials, as well as immediately applicable engineering information. The present progress report describes the program which includes accelerated, laboratory-microbiological tests, as well as the acquisition of data from actual marine exposures. In biochemical oxygen demand-type tests conducted to date polyethylene was not utilized as a carbon source by marine bacteria. Polyvinyl chloride plastics served as a source of energy for the organisms depending on the way in which the materials were plasticized. Five elastomers were utilized by the bacteria. There has been a steady rise in capacitance values for GRS-insulated conductors exposed in sea water and sediment under laboratory conditions for thirteen months. These increases appear due to biological activity on the insulation. The general performance of materials undergoing marine exposure is reported including reference to penetration of a few synthetic materials by marine borers. Brief mention is made of the examination of submarine cable samples from service.

I. INTRODUCTION

As a result of the increasing use of submarine telephone cable, there is a growing demand for information on the performance of organic materials and cable structures under marine conditions. Particularly important is the need for data on the resistance of materials to attack by marine organisms. Although considerable published information exists on the behavior of natural organic materials such as wood, jute, hemp and the like, there is virtually no data on plastics, elastomers, casting resins or similar materials.

About three and a half years ago, the Laboratories initiated a program to determine the resistance to marine biological attack of materials which might find application in submarine cable. The program has two primary objectives: (1) acquisition of fundamental information regarding the biological resistance of a wide range of selected organic materials, and (2) accumulation of immediately applicable engineering information on materials.

II. OUTLINE OF PROGRAM

It seemed evident that marine borers or microorganisms, particularly bacteria, might be expected to be the major agents of deterioration.

Marine borers are mollusks or crustaceans which bore into a material for food or shelter depending on the particular organism involved. Of the crustaceans, the gribble, *Limnoria*, is the most outstanding. Cellulose material such as wood and cordage form its food supply and natural habitat. There are a few references which suggest that members of the genus *Limnoria* have bored into gutta percha. One by Chilton¹ refers to the activity of *Limnoria* in the splice of a submarine cable in about 60 fathoms off New Zealand. Preece² identifies *Limnoria* as the organism responsible for failure of the Holyhead-Dublin cable in 1875. Jona³ states that he frequently found *Limnoria* in cables in the Adriatic Sea. Menzies⁴ points out that no American species is known to occur in depths exceeding 50 feet; however, one species is known to occur off Japan at a depth of 163 fathoms. He suggests that the absence of wood probably limits the distribution of the animals in deep water.

In the bibliography by Clapp and Kenk,⁵ there are ten, separate citations to the attack of submarine cables by molluscan borers belonging to the family *Teredinidae*. In most cases, attack was confined to cellulose constituents such as jute and hemp, although in a few instances mention is made of attack on gutta percha insulation. Although the teredine borers, along with *Limnoria*, are considered to be relatively shallow water organisms, Roch,⁶ in his paper on Mediterranean teredos, refers to *Teredo utriculus* being obtained from depths as great as 3500 meters. There is one reference⁷ to teredo attack of lead-covered submarine cable.

The other important family of boring mollusks is the *Pholadidae*. Members of this family are sometimes referred to as the "burrowing clams" and include rock, shell and wood borers. Some genera, such as *Xylophaga*, are found in water up to 1,000 fathoms or more deep.⁴ Bartsch and Rehder⁸ report the penetration of the lead sheath of a submarine cable by one of the *Martesia*, another genus of the family *Pholadi-*

dae. Members of the same family were reported by Snoke and Richards⁹ to have bored through the lead sheath of a submarine telephone cable.

The bacteria generally are single-celled organisms, a large number of which are heterotrophic, that attack organic matter and use it as a source of carbon or energy. The bacteria play an important part in the biology of the sea, their most important function being to decompose organic material into carbon dioxide, water, ammonia and minerals. The characteristics, distribution and function of the marine bacteria have been described in great detail by ZoBell.¹⁰ Bacteria are found in sea water and sediment from shallow depths to the deepest portions of the sea. During the Danish Galathea Deep-Sea Expedition from 1950 to 1952, bacteria were found in depths as great as 10,280 meters.¹¹ Many of these bacteria have been found to be barophilic²¹, growing exclusively or preferentially at pressures approximating 15,000 psi. ZoBell and Morita¹² have reported experiments performed with these bacteria to determine the effects of high pressure on such factors as viability and enzyme production. Marine bacteria have been found capable of oxidizing rubber products,¹³ as well as a wide variety of gaseous, liquid and solid hydrocarbons.¹⁴ Although evidence to date indicates that among the microorganisms, the bacteria are particularly likely agents of deterioration in the ocean, it is possible that the fungi may also be contributors. Barghoorn and Linder¹⁵ report the physiological behavior and growth on various media of seven species of marine fungi isolated from wood continuously submerged in the sea. Deschamps²⁰ has discussed the role of fungi and bacteria in aiding the attack of wood by marine borers. Also, the occurrence of marine fungi in wood test panels, driftwood and piling in Biscayne Bay has been reported by Myers.¹⁶

A program designed to provide fundamental and engineering data on the susceptibility of organic materials to marine borers and microorganisms in an environment that covers about 70 per cent of the earth's surface could be almost unlimited. The practical parameters which finally were established were based on a number of considerations. Fundamental data on the basic inertness or relative rates of attack by microorganisms could best be determined under controlled laboratory conditions; however, more than one procedure would be needed to determine performance in the environments of water and sediment. Because of the relatively rapid activity of marine borers under natural conditions, and their critical requirements as far as laboratory culture is concerned, it was decided that any borer tests would be conducted in the field. This meant that the natural exposure tests would serve as correlative tests

for the laboratory microbiological portion of the program, as well as a means of evaluating the relative performance of materials in the physical and chemical conditions of the ocean.

The integrated program, shown in Fig. 1, involves a series of three laboratory tests on the one hand, and actual marine exposure on the other. Eventually, more than fifty different materials including plastics, elastomers, natural and synthetic fibers, as well as sections of cable will be tested. The present paper is in the nature of a progress report in which only a portion of the data to be acquired are presented. The results of the program to date will be examined beginning with the laboratory experiments.

III. LABORATORY TESTS

3.1 *Biochemical Oxygen Demand Type Test*

The BOD (biochemical oxygen demand) type of test as applied in this study is really composed of two separate bioassay procedures. In one case, the oxygen consumed by aerobic bacteria is determined, and in the other a metabolic by-product resulting from anaerobic activity is measured. With a few changes, both methods follow those which have been employed by ZoBell¹³ in tests of elastomers and various natural organic materials.

There is one point which should be emphasized regarding both the aerobic and anaerobic procedures used in this accelerated sea water test. It is considered primarily a screening test which provides basic data on

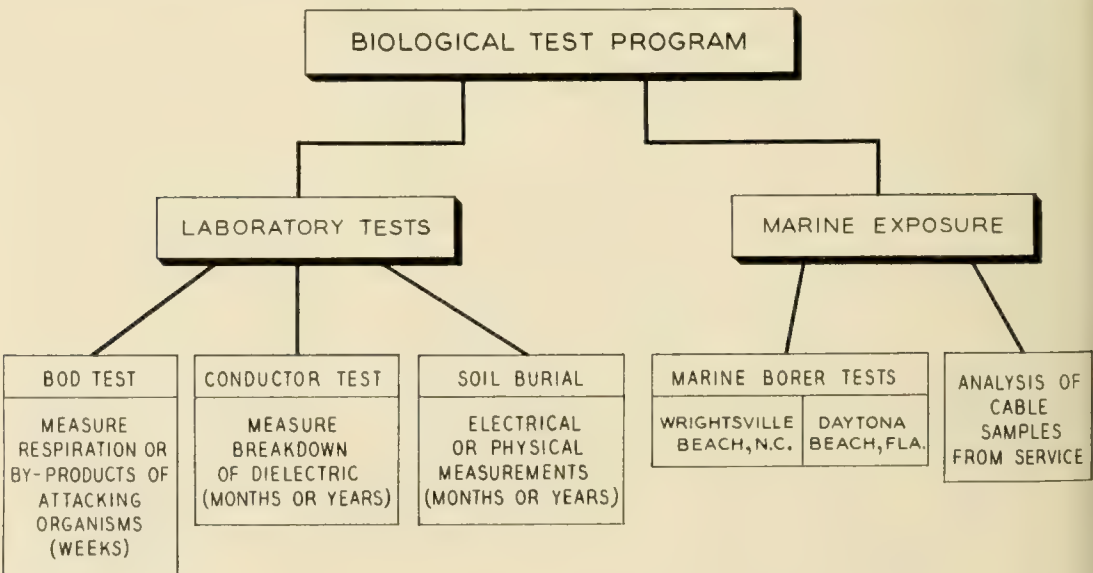


Fig. 1 — Outline of marine biological test program.

the ability of marine bacteria to utilize a compound as a carbon source at the time of test. It will not reflect changes brought about in the material due to prolonged exposure in sea water, or ecological relationships which might make the material more or less susceptible to attack by certain bacteria. The other laboratory tests and natural exposure test will help furnish data on questions involving changes in materials due to long-term marine exposure.

TABLE I — MATERIALS TESTED AGAINST AEROBIC AND ANAEROBIC MARINE BACTERIA

Designation	Type
<i>Polyethylene</i> ¹	
2.0 melt index ²	P5310466
0.2 melt index (Source A)	P5304156
0.2 melt index (Source B)	P5312587
0.2 melt index + 5% butyl rubber + antioxidant	P5308396
0.2 melt index + antioxidant	P5308390
0.7 melt index (high density) nat. + antioxidant	P5503135
0.7 melt index (high density) + carbon black and antioxidant	P5503133
<i>Poly(Vinyl Chloride)</i> ³	
<i>Plasticizer</i>	
Polyester A	BTL 24-54
Di-2-ethylhexyl phthalate (DOP) Shore A 88	BTL 23-54
Tricresyl phosphate (TCP)	BTL 529-53
None (rigid)	P5503087
None ⁴	P5502078
None ⁴	P5502077
Tri-2-ethylhexyl phosphate	P5502081
Nitrile rubber/polyester C	P5502074
Di-2-ethylhexyl phthalate (DOP) Shore A 62	P5502082
Nitrile rubber	P5502076
Polyester E/DOP (BTL 46-55)	P5503115
None (PVC resin)	P5510645
<i>Castings Resins</i>	
Epoxide (cast), unfilled straight epoxy resin cured with amine hardener	
Styrene polyester, silica-filled	
<i>Elastomers</i> ⁵	
GR-S jacket	BTL 54-14
GR-A jacket	BTL 54-18
Butyl jacket	BTL 54-19
Natural rubber jacket	BTL 54-23
Neoprene jacket	BTL 54-164
<i>Jute</i>	

¹ Except where noted polymers are low density grades manufactured by the high pressure process.

² ASTM D1238

³ With the exception of P5510645 all PVC compositions contained typical organo-metallic type stabilizers (such as Ba, Cd, and Pb), fatty acid lubricants in low concentrations, and in some cases small quantities of inorganic fillers.

⁴ Semi-flexible PVC copolymer.

⁵ These compounds all contain, in addition to the basic elastomers, sulfur, accelerators, waxes, processing oils, and reinforcing quantities of carbon black. (54-164 also contains clay.)

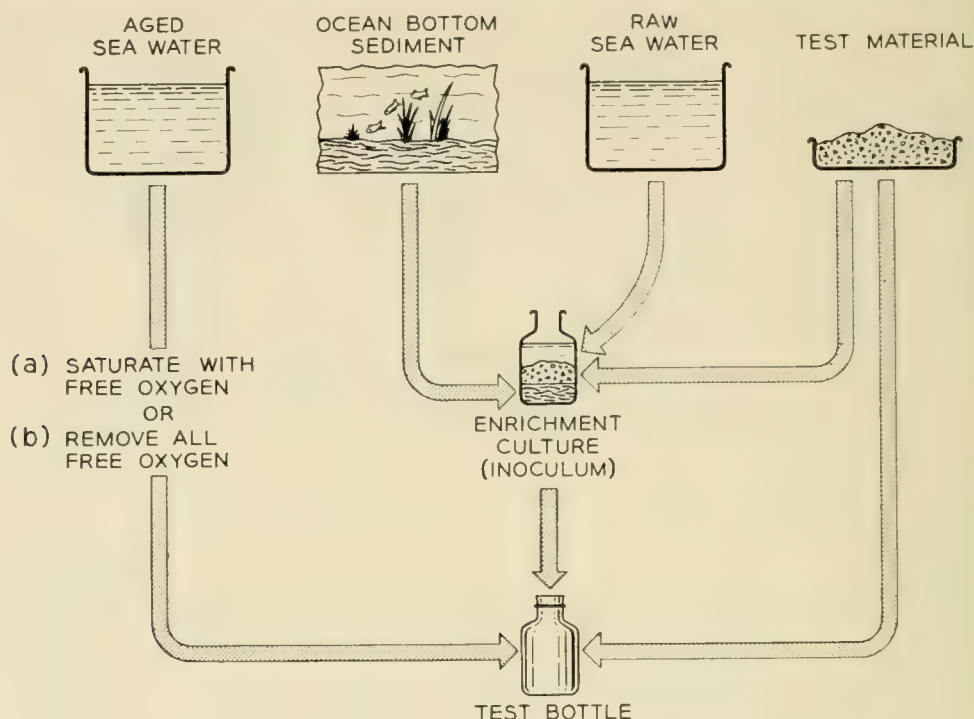


Fig. 2 — Flow chart of biochemical oxygen demand (BOD) type test.

The materials which have been tested thus far by the BOD type procedure include polyethylene, polyvinyl chloride plastics, casting resins, elastomers and jute. The individual materials are listed in Table I. Other plastics and elastomers are still to be tested.

The general features of both the aerobic and anaerobic parts of the test procedure are shown in the flow chart in Fig. 2. Certain features of the test are common to both parts. These features will be described first. The four primary constituents of the test are aged sea water, test material, ocean bottom sediment, and raw (unaltered) sea water. Aged sea water is raw sea water which has been filtered through Millipore filters of 0.5 micron pore size, and then aged in the dark until the biochemical oxygen demand (BOD) is quite low; i.e., until the water contains about 1 ppm of organic matter. This usually requires about eight weeks of aging. In the first tests which were run, materials were finely ground so as to expose a large surface area, and so accelerate attack. However, it soon became apparent in efforts to relate the rate of oxidation to surface area exposed that only crude estimates could be made of the irregular surface areas. Consequently, after the first few tests, thin sheets of material were employed wherever possible so that a measured amount of surface area could be exposed in each case.

The inoculum for the test comes from specially prepared enrichment cultures. Approximately 90 cc of marine sediment is placed in a 250 ml

prescription bottle. About 1 gram of a finely divided test material is also placed in the bottle which is then filled about three-quarters full of raw sea water. To include as heterogeneous a population of marine bacteria as possible, another inoculum is prepared for addition to the enrichment culture. The additional inoculum is made by placing in a vial a small particle of each of seven different sediments furnished by Dr. ZoBell of the Scripps Institute of Oceanography. These sediments are identified in Table II. Following this, one or two drops of liquid are added to the same vial from each of twenty-nine different enrichment cultures which also were provided by Dr. ZoBell. These cultures are identified in Table III. Transfers from eight different cultures of marine sulfate-reducing bacteria are included, and the vial shaken thoroughly. About five drops of pooled inoculum are added to the enrichment culture prepared for each test material. The completed enrichment cultures are incubated at 25° C for a minimum of six weeks prior to use. During the incubation period those bacteria in the culture which are capable of utilizing the test material tend to develop preferentially.

The same enrichment culture is used whether the test procedure is aerobic or anaerobic since both conditions prevail in this type of enrichment culture — aerobic in the water and upper sediment, and anaerobic in the deeper, compacted sediment. From this point on, in describing the method used in the material tests, it is necessary to describe the aerobic and anaerobic procedures separately.

In the aerobic tests, 0.01 per cent ammonium phosphate is added to sufficient sea water (usually about 7 liters) for a given test run. Oxygen is bubbled through the sea water in a carboy for a minimum of sixteen hours at which time the oxygen content of the sea water is about 25 ppm. Since as many as four or more test materials may be included in a test run, the inoculum is prepared by combining in one vial a small amount of liquid from the enrichment culture for each material to be

TABLE II — SOURCES OF SEDIMENTS* USED IN PREPARING ENRICHMENT CULTURES

Ref. No.	Source
XG 17-4 (surface).....	Gulf of Mex., Rockport, Texas
5403-1 (surface).....	Gulf of Mex., Miss. Delta
XS-384 (surface).....	Gulf of Mex., Rockport, Texas
5402-7 (0-5 cm).....	Gulf of Mex., Miss. Delta
56:180 (4520 fathoms).....	Pacific, 7° 22.2'N, 127° 17'W
56:184 (2650 fathoms).....	Pacific, 19° 02'N, 174° 58'W
56:177 (4550 fathoms).....	Pacific, 7° 03.8'N, 126° 24.3'W

* Obtained from Dr. C. ZoBell, Scripps Institute of Oceanography.

included in the run. Once in the vial, the inoculum is shaken and added to the aged sea water at the rate of 1 ml per 10 liters of medium. This amount of inoculum was calculated to give the maximum number of bacteria, consistent with a minimum addition of organic matter. The carboy is then placed under slight, positive oxygen pressure.

The test is run in 60 ml glass-stoppered bottles. A small amount of test material is placed in each bottle. At the outset of the experiments, when ground material was used, this amounted to 0.05 gram, or a surface area of 4 to 45 sq cm, depending on the material. Later, when thin sheets of about 4 mils thickness were used, the samples were cut to a size of 2.54 cm square. The samples are placed in the bottles the night before, and enough aged sea water added to permit surface wetting. With many materials this seems to result in less accumulation of air bubbles on the surfaces of the materials during subsequent filling with the medium.

TABLE III — ENRICHMENT CULTURES* USED AS SUPPLEMENTARY SOURCES OF INOCULUM IN PREPARATION OF ENRICHMENT CULTURES FOR CURRENT PROGRAM

Ref. No.	Description
34-134	rubber in distilled water
34-134	anthracene in sea water
34-134	sewage outfall, rubber in sea water
34-134	mixed hydrocarbons in sea water
34-134	garden soil, rubber in sea water
34-132	Athabaska tar sand, hydrocarbon-oxidizing bacteria and sulfate-reducers in sea water
34-134	tricresol in sea water
34-134	mixed hydrocarbons in sea water
25-143	0.10% phenol in sea water
34-134	0.25% phenol in sea water
34-134	cork in sea water
34-134	Shell oil No. 10 in sea water
25-141	lignin in sea water
34-134	sewage outfall, rubber in sea water
34-134	sawdust and mud in sea water
34-134	garden soil, rubber in sea water
34-134	rubber in tap water
34-134	mixed crude oil in sea water
34-134	kerosene in sea water
34-134	paraffin in sea water
34-134	rubber in tap water
—	Athabaska tar sand, mixed crude oil in sea water
—	thiokol in sea water
—	neoprene in sea water
—	cellulose acetate in sea water
—	butadiene (Buna A) in sea water
—	pooled aerobic hydrocarbon-oxidizing bacteria in sea water
—	crude coal tar in sea water
—	shellac in sea water

* Obtained from Dr. C. ZoBell, Scripps Institute of Oceanography.

Of course, air bubbles would be a source of error in later oxygen determinations.

Usually, sufficient test bottles are made up to provide duplicates for analysis after each period of incubation. Oxygen pressure, maintained on the sea water in the carboy, assures no loss of oxygen from the medium and forces it through tubing into the test bottles. Since incubation periods of 0, 1, 2, 4 and 8 weeks are used as a general guide, and two test bottles must be sacrificed for analysis after each interval, ten test bottles are used for each material. One set of ten control bottles containing only inoculated, aged sea water suffices for a test run, as long as the bottles for materials and controls are made up from the same batch of sea water and incubated at the same time. Incubation is carried out in the dark in a constant temperature water bath maintained at 20° ± 0.5°C. Incubation in the water bath minimizes the fluctuation in oxygen content of the sea water which might be encountered as the result of "breathing" of the bottles in atmospheric incubation. After the various incubation periods, the free oxygen content of the sea water in the bottles is determined by a modified Winkler procedure.

In the anaerobic portion of the test, the procedure is essentially the same as for the aerobic part, the only differences being in the preparation and handling of the sea water medium, the incubation times, and the analytical method. Of course, with the anaerobic bacteria it is necessary to remove all free oxygen from the sea water medium if the organisms are to function. Consequently, instead of bubbling oxygen through the medium, the sea water is boiled for ten minutes and placed hot in a

TABLE IV — OXYGEN CONSUMPTION BY MARINE BACTERIA
IN BOD TEST WITH POLYETHYLENE AS THE
ONLY SOURCE OF ORGANIC CARBON

Test Material ²	O ₂ Consumption After Weeks of Incubation			
	1	2	4	8
	ppm	ppm	ppm	ppm
2.0 melt index ³	2.2	3.7	6.0	10.2
0.2 melt index (Source A)	0.8	1.6	2.4	10.9
0.2 melt index (Source B)	1.4	3.4	5.2	10.8
0.2 melt index + antioxidant	0.9	3.0	6.2	8.5
0.2 melt index + 5% butyl rubber + antiox.	0.2	3.8	5.8	— ¹
0.7 melt index (High Density) + antiox.	0.9	4.3	7.4	9.3
Controls (inoculated sea water)	1.5	5.3	8.3	11.2

¹ Samples accidentally destroyed.
² Except where noted polymers are low density grades manufactured by the high pressure process.
³ ASTM D1238

TABLE V — OXYGEN CONSUMPTION BY MARINE BACTERIA IN BOD TEST WITH POLY(VINYL CHLORIDE) PLASTICS, EPOXIDE CASTING RESIN OR JUTE AS ONLY SOURCES OF ORGANIC CARBON

Test Material	O ₂ Consumption After Weeks of Incubation			
	1	2	4	8
	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>
PVC — no plasticizer (rigid).....	11.1	12.9	11.6	18.7
PVC — tricresyl phosphate (TCP).....	9.5	13.2	21.6	22.2
PVC — di-2-ethylhexyl phthalate (DOP) Shore A 88.....	9.1	13.4	19.7	20.7
PVC — polyester A.....	19.3	22.2	*	*
Epoxide casting resin.....	—	4.1	5.1	4.2
Jute.....	10.0	15.0	16.5	*
Controls (inoculated sea water).....	6.8	6.8	7.7	7.7

* All free O₂ in sea water consumed.

carboy containing 0.01 per cent ammonium phosphate. Nitrogen is introduced into the carboy immediately. When the sea water is cool, inoculum, which is prepared as described for the aerobic procedure, is added and additional nitrogen pressure placed on the carboy for filling the test bottles. Since anaerobic activity is usually slower than aerobic, the time in test is increased. Analysis for hydrogen sulfide in the sea water is carried out at 0, 4, 8, 12 and 16 weeks. Since the sulfate-reducing bacteria are ubiquitous anaerobic marine species, the hydrogen sulfide produced by them in the course of breaking down organic material is used as an indicator of their activity. The sulfide in the sea water is determined volumetrically according to the method described in the Official

TABLE VI — OXYGEN CONSUMPTION BY MARINE BACTERIA IN BOD TEST WITH POLY(VINYL CHLORIDE) PLASTICS AS THE ONLY SOURCE OF ORGANIC CARBON

Plasticizer	O ₂ Consumption After Weeks of Incubation			
	1	2	4	8
	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>
Nitrile rubber/polyester C.....	10.3	12.9	21.4	*
Nitrile rubber.....	9.2	12.3	18.7	*
None ¹	3.7	4.2	6.5	10.5
None ¹	4.0	5.5	8.4	11.0
Tri-2-ethylhexyl phosphate.....	11.7	14.4	23.1	*
Di-2-ethylhexyl phthalate (DOP) Shore A 62.....	6.4	8.4	11.5	*
Polyester E/DOP (BTL 46-55).....	*			
Controls (inoculated sea water).....	3.5	4.5	7.1	9.7

* All free O₂ in sea water consumed.

¹ Semi-flexible PVC copolymer.

TABLE VII — OXYGEN CONSUMPTION BY MARINE BACTERIA IN BOD TEST WITH POLYETHYLENE, POLYESTER CASTING RESIN OR POLY-(VINYL CHLORIDE) RESIN AS ONLY SOURCE OF ORGANIC CARBON

Test Material	O ₂ Consumption After Weeks of Incubation			
	1	2	4	8
	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>
Polyethylene 0.7 melt index (High Dens.) Nat. + antioxidant.....	3.1	3.3	6.1	6.5
Polyethylene 0.7 melt index (High Dens.) Blk.....	2.9	4.1	7.1	7.0
Styrene polyester, silica-filled.....	5.5	7.0	9.3	12.1
Poly(Vinyl Chloride) resin.....	4.2	4.1	7.3	7.2
Controls (inoculated sea water).....	2.5	3.8	6.7	6.7

and Tentative Methods of Analysis of the Association of Agricultural Chemists.

The results of the aerobic test procedure are presented in Tables IV to VIII inclusive. In these tables the oxygen consumption values obtained with materials which went through the same test run are included in the same table. The materials included in Tables IV and V, with the exception of jute, were exposed in finely ground or shaved form. Since it was not possible to obtain reliable estimates of the surface areas exposed to attack in this case, the oxygen consumption values in these tables are not directly comparable with respect to rate. The data serve the important basic purpose of indicating whether these materials can serve as a source of energy for the bacteria. However, data in Tables VI to VIII inclusive are based on the use of equally thin sheets of material with about 12.9 sq cm of surface area exposed to attack. Two exceptions

TABLE VIII — OXYGEN CONSUMPTION BY MARINE BACTERIA IN BOD TEST WITH ELASTOMERS AS THE ONLY SOURCE OF ORGANIC CARBON

Elastomer	O ₂ Consumption After Days of Incubation			
	3	7	14	28
	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>
GR-S jacket (54-14)	15.8	*		
GR-A jacket (54-18)	6.1	10.9	23.5	*
Butyl jacket (54-19)	13.9	*		
Natural rubber jacket (54-23)	14.1	*		
Neoprene jacket (54-164)	1.9	4.1	10.3	*
Controls (inoculated sea water).....	0.0	-0.4	0.0	

* All free O₂ in sea water consumed.

to this are the polyvinyl chloride resin and styrene polyester in Table VII which could not be prepared in sheet form.

Results for polyethylene are described in Tables IV and VII. The oxygen consumption values in these tables are almost identical to the control values obtained using only inoculated sea water. There is no evidence in any of these tests of polyethylene being utilized as a source of carbon by the bacteria. In fact, in Table IV, all of the eight week values for polyethylene are slightly below those for inoculated sea water alone.

The results with the polyvinyl chloride plastics vary according to the manner in which the compounds are plasticized. The data are contained in Tables V, VI and VII. First, as may be noted in Table VII, there is no attack on the polyvinyl chloride resin. This indicates that the susceptibility of these plastics can be attributed to materials added in compounding. Every polyvinyl chloride plastic tested shows some evidence of attack; (distinct oxygen consumption above the control rate), except the semi-flexible copolymers which contain no added external plasticizer. In these compounds, acrylates are employed as copolymers. The most severe attack occurred on the plastic in Table VI which con-

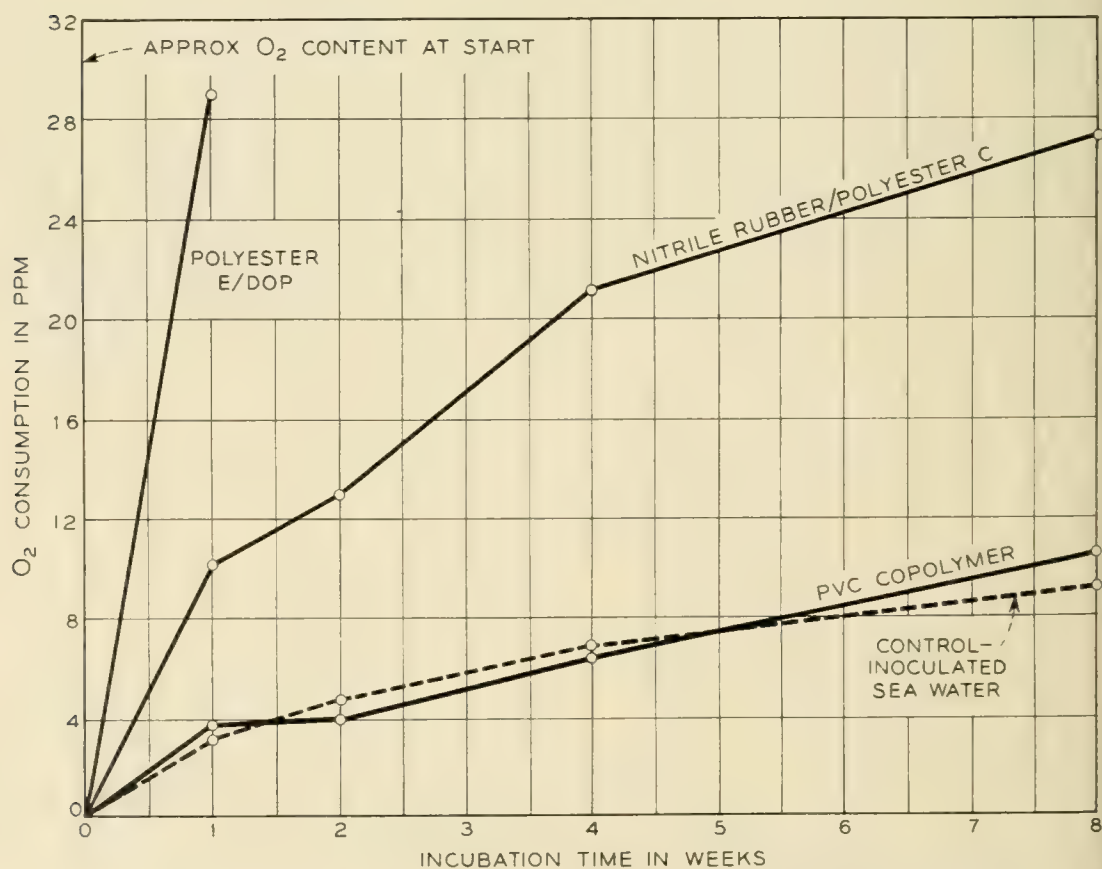


Fig. 3 — Examples of O_2 consumption by marine bacteria in BOD test with Poly(Vinyl Chloride) plastics as carbon source.

tains a combination of polyester E and di-2-ethylhexyl phthalate (DOP), and one in Table V plasticized with polyester A. These polyesters are fatty acid-type compounds. Typical oxygen consumption values for three different polyvinyl chloride plastics, representing different rates of utilization by the bacteria, are plotted in Fig. 3. As noted in Table I, very low concentrations of organo-metallic stabilizers and fatty acid lubricants were in all the compounds tested except the resin alone. However, of the three materials which contained no added external plasticizer only the rigid plastic is utilized. This material contained about 8 to 10 times as much fatty acid lubricant as the other two compounds.

Two casting resins were tested, one an epoxide (Table V), and the other a silica-filled styrene polyester (Table VII). Under the conditions of this test, the epoxide resin is not utilized by the organisms. In the case of the styrene polyester, results are less conclusive. After eight weeks, an oxygen consumption value 5.4 ppm higher than that for the controls suggests the possibility of attack. Additional tests are planned with this material to obtain more data on which to base a final decision.

As might be expected, the jute fibers are quite susceptible to attack; all oxygen was consumed from the test medium between the fourth and eighth week (Table V). The fact that results in the same test run with the polyvinyl chloride compound plasticized with polyester A show that all oxygen was consumed from the test medium in 17 days does not mean that this latter compound is more susceptible to attack than jute. In the jute, bacterial attack is necessarily restricted to a progressive surface attack, but with the polyvinyl chloride compound, leaching of the sus-

TABLE IX — HYDROGEN SULFIDE PRODUCTION BY MARINE BACTERIA IN ANAEROBIC SEA WATER TEST WITH POLYETHYLENE AS THE ONLY SOURCE OF ORGANIC CARBON

Test Material ¹	H ₂ S Production After Weeks of Incubation			
	4	8	12	16
	ppm	ppm	ppm	ppm
2.0 melt index ²	0.22	0.32	0.22	— ³
0.2 melt index (Source A)	0.22	0.29	0.45	— ³
0.2 melt index (Source B)	0.22	0.32	0.26	0.38
0.2 melt index + antioxidant	0.22	0.64	0.35	0.58
0.2 melt index + 5% butyl rubber + antiox.	0.22	0.32	1.31	0.24
0.7 melt index (High Density) + antiox.	0.22	0.22	0.29	0.45
Controls (inoculated sea water)	0.10	0.64	0.70	0.58

¹ Except where noted polymers are low density grades manufactured by the high pressure process.

² ASTM D1238

³ Insufficient samples

ceptible plasticizer into the sea water medium might greatly accelerate utilization of that material and be reflected in rapid oxygen consumption.

Five different elastomers have been evaluated by the BOD test to date. The results with aerobic bacteria are presented in Table VIII. First, it is apparent that all of the elastomers tested can serve as a source of carbon for the bacteria. As may be noted in the table, GR-S jacket (54-14), butyl jacket (54-19) and natural rubber (54-23) are oxidized at about the same rate—all oxygen being consumed from the test medium between the third and seventh day analyses. GR-A (54-18) and neoprene jacket (53-164) are more resistant than the other three elastomers in the test. During the fourteen-day test period, approximately twice as much oxygen was consumed in the case of the GR-A as with the neoprene.

The results of anaerobic bacterial activity, as reflected by analyses for hydrogen sulfide in the sea water medium, are contained in Tables IX to XII, inclusive. As with the results of the aerobic test, materials in a given test run are included in the same table. No polyethylene is utilized as a source of carbon by the sulfate-reducing bacteria. In no case is the production of hydrogen sulfide, with different polyethylenes

TABLE X — HYDROGEN SULFIDE PRODUCTION BY MARINE BACTERIA IN ANAEROBIC SEA WATER TEST WITH POLY(VINYL CHLORIDE) PLASTICS, EPOXIDE CASTING RESIN, OR JUTE AS ONLY SOURCES OF ORGANIC CARBON

Test Material	H ₂ S Production After Weeks of Incubation			
	4	8	12	16
	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>
Poly(Vinyl Chloride) Plastics				
No plasticizer (rigid).....	1.90	3.50	4.20	5.60
Tricresyl phosphate (TCP).....	0.22	0.22	0.64	0.38
Di-2-ethylhexyl phthalate (DOP) Shore A 88.....	0.26	0.26	0.26	0.61
Polyester A.....	2.60	38.40	38.40	47.70
Nitrile rubber/polyester C.....	1.50	7.00	19.98	18.60
Nitrile rubber.....	4.10	9.60	12.40	11.50
No plasticizer ¹	0.22	0.22	0.90	0.51
No plasticizer ¹	0.32	0.64	1.89	1.02
Tri-2-ethylhexyl phosphate.....	0.26	0.96	1.86	0.99
Di-2-ethylhexyl phthalate (DOP) Shore A 62.....	0.22	0.22	1.09	0.58
Polyester E/DOP (BTL 46-55).....	12.20	61.40	94.10	82.90
Epoxide casting resin.....	0.16	0.64	0.86	0.48
Jute.....	1.90	13.40	38.60	52.20
Controls (inoculated sea water).....	0.10	0.64	0.70	0.58

¹ Semi-flexible PVC copolymer

TABLE XI — HYDROGEN SULFIDE PRODUCTION BY MARINE BACTERIA
IN ANAEROBIC SEA WATER TEST WITH POLYETHYLENE,
POLYESTER CASTING RESIN OR POLY(VINYL CHLORIDE)
RESIN AS ONLY SOURCES OF CARBON

Test Material	H ₂ S Production After Weeks of Incubation			
	4	8	12	16
	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>
Polyethylene — 0.7 melt index (High Dens.) Nat. + antioxidant	0.35	0.58	0.90	1.12
Polyethylene — 0.7 melt index (High Dens.) Blk. + antioxidant	0.26	0.48	0.38	0.74
Silica-filled styrene polyester.....	0.83	0.83	1.02	1.02
Poly(Vinyl Chloride) resin.....	0.48	0.58	0.58	0.58
Controls (inoculated sea water).....	0.45	0.86	0.91	0.91

as the test material (Tables IX and XI), significantly greater than in the control bottles. In fact, in most cases it is actually less than that for the controls.

Four polyvinyl chloride plastics appear to have served as a source of carbon for the anaerobic organisms. In order of decreasing susceptibility they are the compounds plasticized with (1) polyester E/DOP, (2) polyester A, (3) nitrile rubber/polyester C, and (4) no plasticizer (rigid). Just as in the case of the aerobic procedure, the plastic plasticized with polyester E/DOP is used much more rapidly than any of the other polyvinyl chloride compounds. There is no evidence of attack on polyvinyl chloride resin, again indicating that the attack is on the plasticizers, not the polyvinyl chloride itself. In this regard, it should be pointed out again that although the polyvinyl chloride compound listed as “no plasticizer (rigid)” in the tables and text does not contain an external plas-

TABLE XII — HYDROGEN SULFIDE PRODUCTION BY MARINE BACTERIA
IN ANAEROBIC SEA WATER TEST WITH ELASTOMERS AS THE
ONLY SOURCES OF ORGANIC CARBON

Elastomer	H ₂ S Production after Weeks of Incubation			
	4	8	12	16
	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>	<i>ppm</i>
GR-S jacket (54-14)	0.66	0.77	0.56	0.75
GR-A jacket (54-18)	0.96	0.95	0.72	0.87
Butyl jacket (54-19)	1.52	6.64	8.03	9.65
Natural rubber jacket (54-23)	1.62	2.48	3.58	3.74
Neoprene jacket (54-164)	1.06	1.20	0.99	0.96
Controls (inoculated sea water).....	0.37	0.43	0.43	0.42

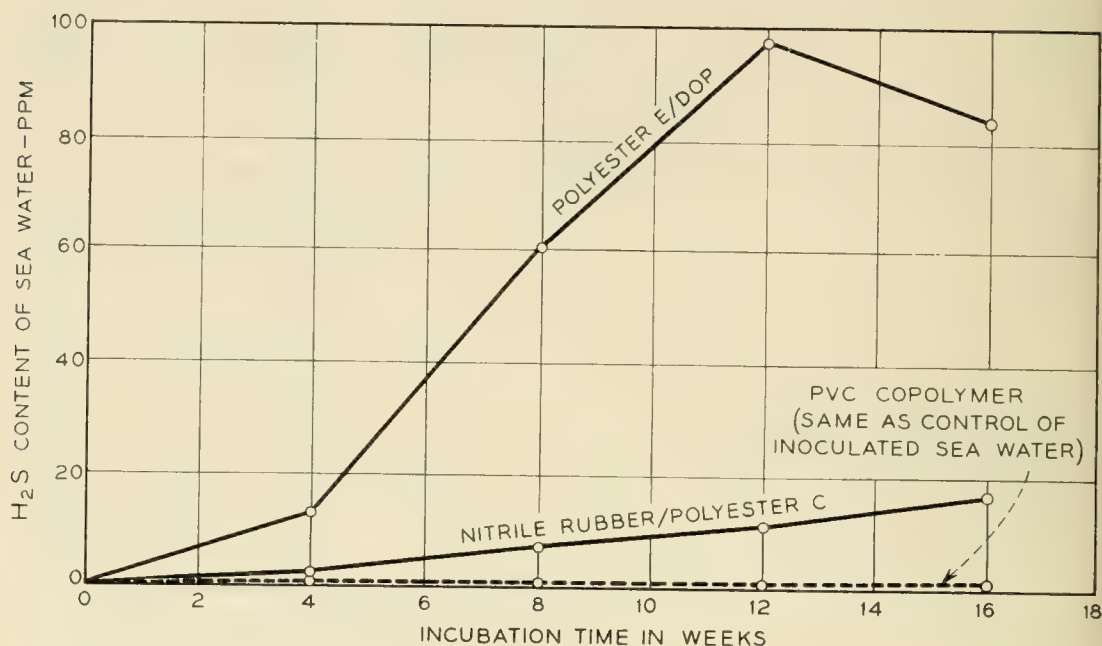


Fig. 4 — Examples of H_2S production by marine bacteria in anaerobic sea water test with Poly(Vinyl Chloride) plastics as carbon source.

ticizer, fatty acid lubricant probably serves as a source of nutrient. For comparative purposes, the examples of hydrogen sulfide production plotted in Fig. 4 are for the same plastics included in Fig. 3 which relates to the data from the aerobic procedure.

Jute is attacked by the anaerobic bacteria just as it is utilized by the aerobic organisms. However, neither the epoxide casting resin (Table X) or the polyester casting resin (Table XI) seem to serve as a source of carbon.

The results of the anaerobic test with the elastomers are presented in Table XII. It is interesting to note that early attack occurs on the natural and butyl rubber jackets, but that none of the other elastomers is utilized by the organism. It is somewhat surprising that attack on GR-S did not progress at about the same rate as on natural and butyl rubber.

The data which have been obtained in the aerobic and anaerobic parts of the BOD-type test are summarized in Table XIII. There is one outstanding fact about the data — no material was utilized by the anaerobic bacteria which was not utilized also by the aerobic organisms. Under the conditions of the test, however, materials did serve as a carbon source for aerobic bacteria and not for the anaerobes.

3.2 Conductor Test

It is apparent that the BOD test provides considerable fundamental information on the ability of halophilic bacteria to utilize organic ma-

terials as a carbon source in sea water. There is little or no opportunity for ecological factors to come into play, however, particularly with regard to marine sediment. In the conductor test, sea water and marine sediment form a part of the test environment, and the test is run over a much longer period of time, thus encouraging more natural and dynamic organism associations. Likewise, the natural relationship between material and environment is simulated more closely than it is in the more accelerated test. In these respects, the conductor test is intermediate to the BOD-type test and natural marine exposure.

The material to be tested is coated on a conductor to provide about 10 mils of insulation. A standard coil of this insulated conductor is then exposed in a 16-ounce bottle so that half of the coil is in marine sediment, and half is in sea water. The ends of the coil are brought through holes in the bottle cap and attached to terminals in the cap. The general features of the test setup are shown in Fig. 5. The bottle is incubated at 20°C. Capacitance and conductance measurements, taken monthly, indicate any change in the insulation. Some conductors are placed in sterile sea water and sediment to serve as controls. This type of test can be continued for months or years if necessary.

Most of the conductor tests are now being initiated. Two materials, however, GR-S and a rigid polyvinyl chloride, have been under study

TABLE XIII—SUMMARY OF MATERIALS UTILIZED AS SOURCE OF CARBON BY AEROBIC OR ANAEROBIC MARINE BACTERIA IN BOD-TYPE TEST

Utilized as Source of Carbon by	
Aerobic Bacteria	Anaerobic Bacteria
PVC — no plasticizer (rigid)	PVC — no plasticizer (rigid)
PVC — tricresyl phosphate (TCP)	
PVC — di-2-ethylhexyl phthalate (DOP) Shore A88	
PVC — polyester A	PVC — polyester A
PVC — nitrile rubber/polyester C	PVC — nitrile rubber/polyester C
PVC — nitrile rubber	PVC — nitrile rubber
PVC — tri-2-ethylhexyl phosphate	
PVC — di-2-ethylhexyl phthalate (DOP) Shore A 62	
PVC — polyester E/DOP (BTL 46-55)	PVC — polyester E/DOP (BTL 46-55)
Styrene polyester	
GR-S jacket (54-14)	
GR-A jacket (54-18)	
Butyl jacket (54-19)	Butyl jacket (54-19)
Natural rubber jacket (54-23)	Natural rubber jacket (54-23)
Neoprene jacket (54-164)	
Jute	Jute

for several months in a "dry" run to establish the biological procedure, as well as the techniques of measurement which are to be employed. The capacitance and conductance data which have been obtained on these two materials to date are presented in Figs. 6 and 7, respectively. In the case of the test samples, each point represents the average value for four test coils, but for the controls each point represents only one coil.

With the GR-S test samples, there is a sharp rise in the capacitance values between the second and third month, amounting to about $110\ \mu\text{f}$. Thereafter, the rise in the curve continues, the slope decreasing somewhat at about the eighth-month point. Over the entire test period, the capacitance ranged from $632\ \mu\text{f}$ at the start to $850\ \mu\text{f}$ after 13 months



Fig. 5 — General setup of conductor test showing a coil half in sediment and half in sea water.

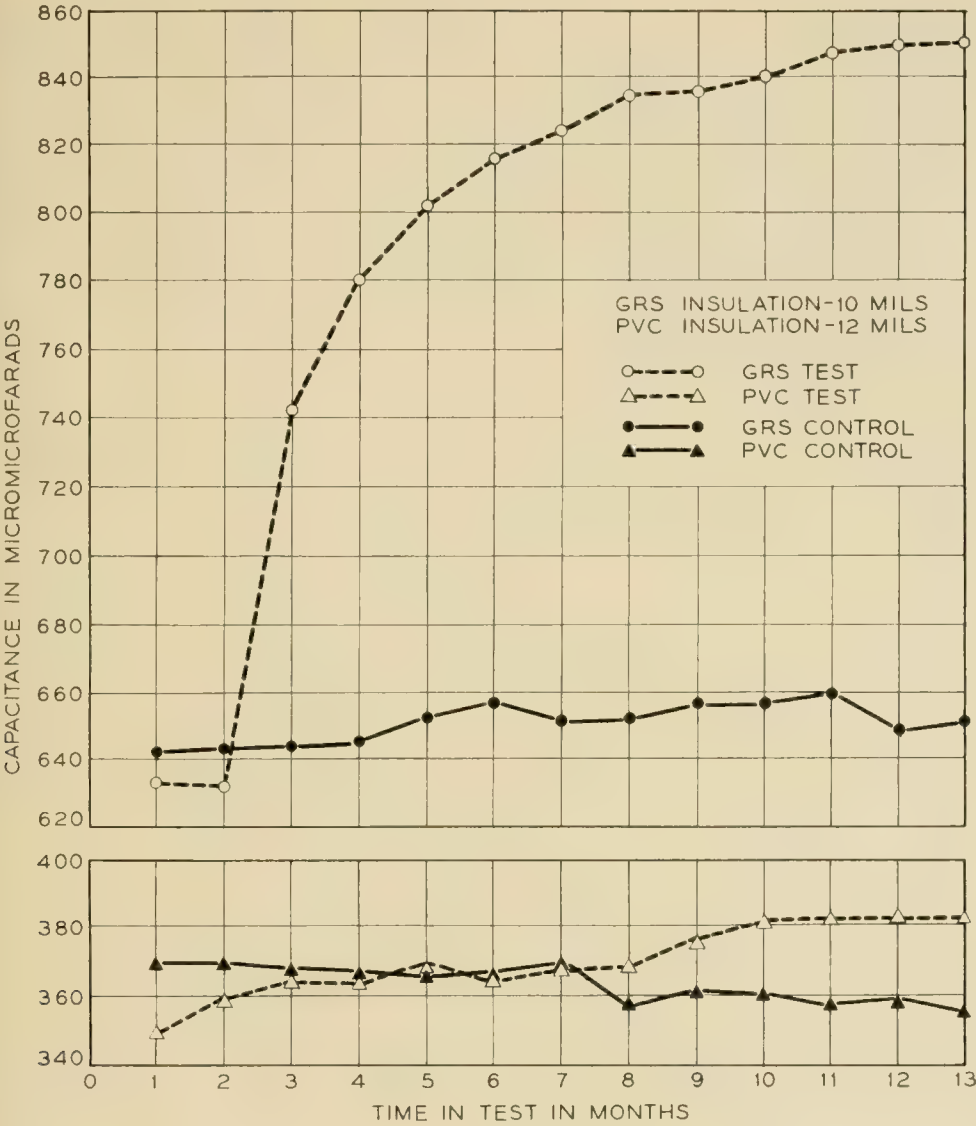


Fig. 6 — Capacitance changes resulting from exposure of GR-S (51-92) and Poly(Vinyl Chloride) (BTL 172-54) insulated conductors in sea water and sediment.

exposure — a total change of 218 $\mu\mu\text{f}$. If it is assumed that the insulating materials were removed equally along the length of the coil, it can be computed that this change in capacitance represents a loss of 8.1 mils of insulation. The following formula is used to arrive at this figure:

$$D = d \left(\frac{D_0}{d} \right)^{C_0/C}$$

- where D = present diameter in mils,
 D_0 = original diameter in mils,
 d = diameter of wire in mils,
 C_0 = original capacitance in $\mu\mu\text{f}$ (start of test),
 C = present capacitance in $\mu\mu\text{f}$.

In contrast, capacitance values for the controls have remained essentially unchanged.

There has been no substantial increase in the capacitance of the polyvinyl chloride-insulated conductors although there is some evidence of an upward trend in the data for the coils in the biologically active environment. There was a rise in capacitance of about $15\ \mu\mu\text{f}$ between the one and three month period, and a similar rise between the eighth and tenth month. Future measurements should indicate whether any biological attack is occurring.

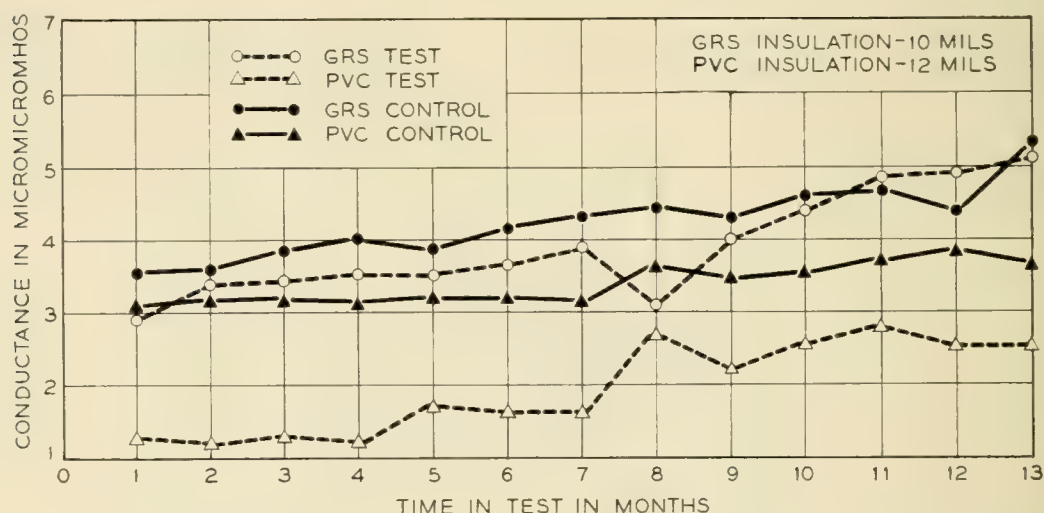


Fig. 7 — Conductance changes resulting from exposure of GR-S (51-92) and Poly(Vinyl Chloride) insulated conductors in sea water and sediment.

As may be noted from the conductance data for both materials in Fig. 7, the patterns of the curves for the test and control samples are essentially the same. In all cases there has been a slight increase in conductance over the 13 month test period. For GR-S it has increased from about 1.5 to $2.0\ \mu\mu\text{mhos}$ and for the polyvinyl chloride from about 0.75 to $1.5\ \mu\mu\text{mhos}$.

In a test of this kind the rise in capacitance, such as that which occurs with the GR-S, without any marked corresponding increase in conductance suggests that the insulation is being modified by the attack, rather than actually removed as assumed previously. It also suggests that the action is most likely due to bacteria rather than fungi which might be expected to penetrate the insulation directly to the conductor and so have a more pronounced effect on conductance. When the test is terminated, it is hoped that these insulated conductors can be run through a capacitance and conductance monitor to locate the specific points of deterioration, and to determine the extent and type of attack.

IV. SOIL BURIAL

Since this phase of the test program is yet to be started no extended coverage is possible in this paper, except to point out the reason for its inclusion. There is some evidence in the results of the marine exposure tests at the Laboratories that the general order of susceptibility of materials to marine microorganisms is the same as it is to terrestrial microorganisms. This observation has also been supported in discussions with some other investigators. The current program at the Laboratories offers an excellent opportunity to compare the performance of a wide range of materials in the two environments. If a correlation pattern can be established, considerably more data in the literature can be brought to bear on the problem. Perhaps at a later date it will be possible to present data comparing material performance in the laboratory soil-burial test and marine-type tests.

V. MARINE EXPOSURE

5.1 *Marine Borer Tests*

The Laboratories, in cooperation with the William F. Clapp Laboratories, Inc., Duxbury, Massachusetts, is conducting the marine borer tests. This phase of the program was initiated in 1954 and involves the natural exposure of specimens at two locations — Wrightsville Beach, North Carolina, and Daytona Beach, Florida. These tests are aimed primarily at obtaining information on marine borer attack; however, the samples are exposed in such a way that information is obtained on microbiological activity as well. In addition, valuable data is obtained on the purely physical and chemical effects of the environment on the materials.

Wrightsville Beach and Daytona Beach were selected as test sites because of the severe and diversified borer activity present in the two areas. At present, more than fifty different materials are exposed at the two locations. Represented are plastics, elastomers and natural organic materials. All of the materials which have been put through the BOD-type test, or are still to be included, are represented in the marine borer portion of the test program. Where possible, test specimens are made in solid rod or tube form about one inch in diameter and three feet long to simulate cable shape. In the case of fibers and tapes, samples are wrapped on $\frac{3}{4}$ -inch diameter Lucite rods 3 feet long. These rods are assembled in racks of about 26 rods each. An untreated, southern pine two by four, susceptible to borer attack, is fitted around the samples at midpoint, where it functions as a bait piece to lead the organisms into direct contact with each test rod. Of course, where there is no bait piece

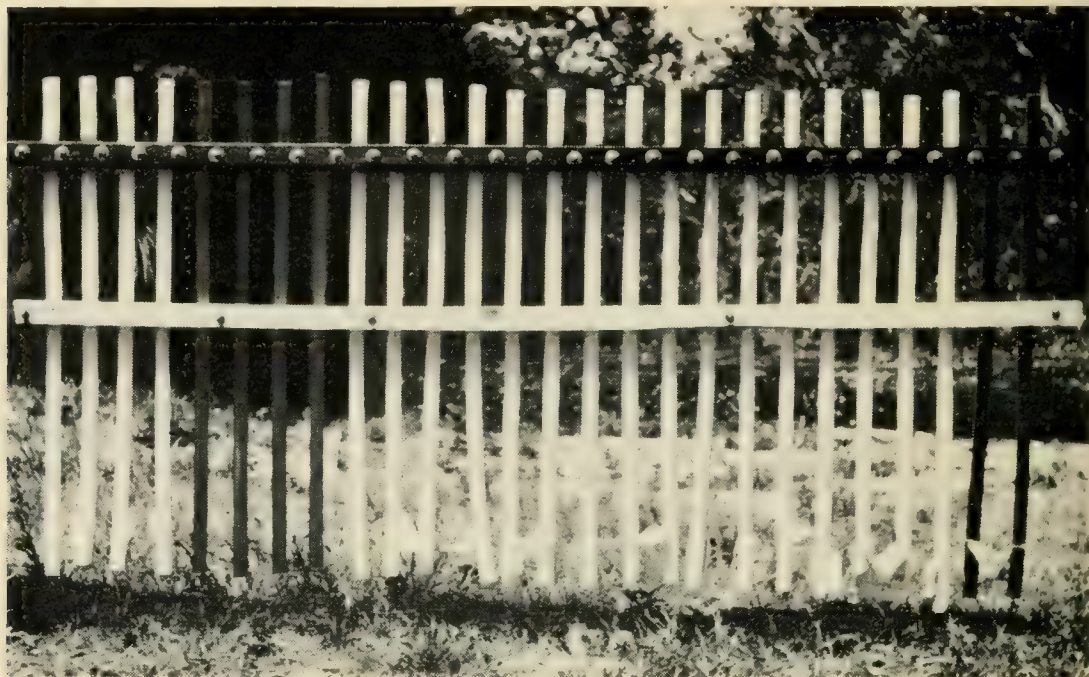


Fig. 8 — A test rack used in marine borer test prior to exposure. Note bait piece of untreated wood fastened across middle of test rods.

it is possible to determine whether the organisms can attack the samples directly from the water. Fig. 8 is a photograph of one of the racks prior to exposure in the sea. The lower 10 inches of the rods are embedded in the bottom sediment where bacterial action is relatively high. Thus, each sample is subjected to water exposure and possible borer attack through the transition zone from water to sediment and into the generally anaerobic conditions of the sediment.

Due to the short time that these tests have been in progress, it is impossible to draw extensive conclusions, particularly with regard to microbiological activity. Long exposure times may bring about physical or chemical changes in a material which may render it more, or less, susceptible to attack. However, until more detailed data are available, some interesting preliminary examples of biological activity can be cited which may be of some interest and serve to illustrate the kind of information which is steadily being acquired.

With but two exceptions, there has been no direct penetration by borers, or microbiological deterioration, of any of the plastics. Polymono-chlor-trifluoroethylene in the form of a 0.0035 inch-thick tape wrapped on a $\frac{3}{4}$ -inch diameter Lucite rod for exposure, was penetrated at one point by a pholad. This attack occurred after three years of exposure at Daytona Beach. Apparently the mollusk bored through an accumulation of calcareous fouling and then progressed through the plastic into the

Lucite rod. In another instance, after $2\frac{1}{2}$ years of exposure at Wrightsville Beach, there was penetration of a silicone rubber test rod at a single point by a pholad. In this case, the test sample was a solid rod of the elastomer one inch in diameter. Attack occurred on the cross-sectional face of the mud end of the rod. Penetration was to a depth of about 4 mm, and the dimensions of the hole at the point of entry were 1.5 by 2.0 mm. Although these examples serve to demonstrate the ability of pholads to bore into these materials, it should be emphasized again that attack has been confined to single points and is not general on these materials.

The physical relationship of one material to another can be very important as far as borer attack is concerned. A piece of one of the Lucite rods on which jute roving was wrapped for exposure is shown in Fig. 9. The holes in the rod resulting from penetration by pholads are readily apparent. One of the organisms may be seen protruding from the left-



Fig. 9 — Section of Lucite rod showing penetration by pholads. One of the mollusks can be seen protruding from the left-hand side of the rod. Original magnification 2X.

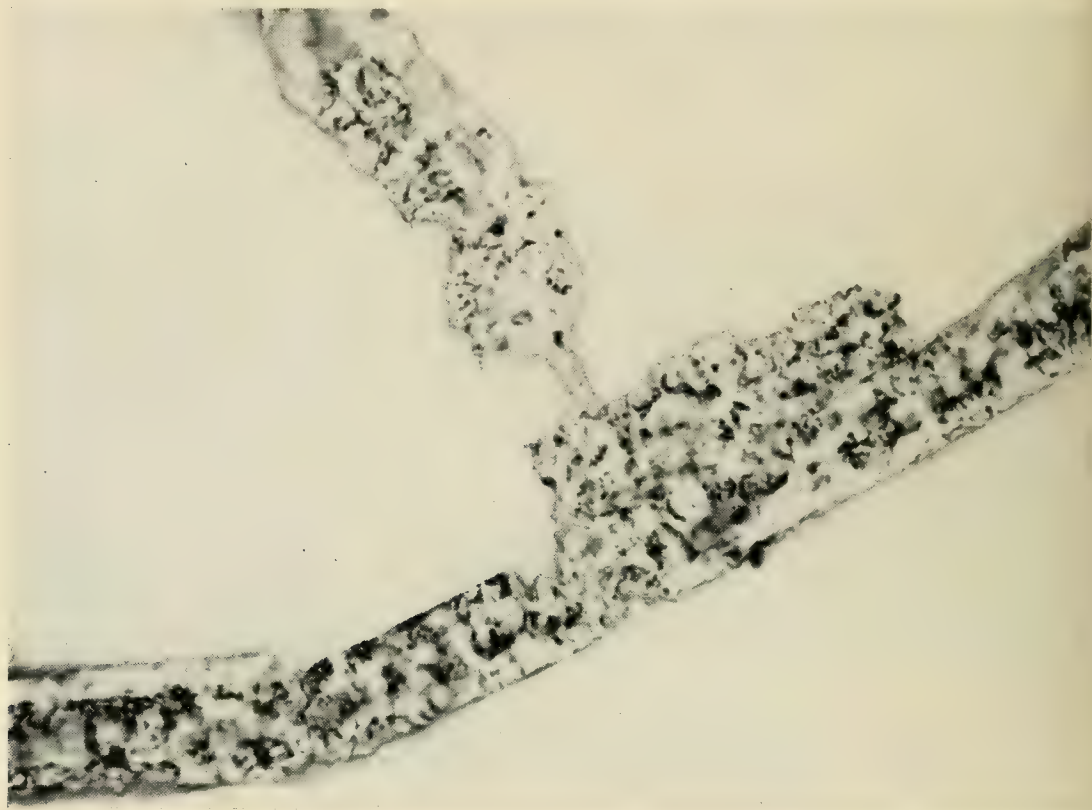


Fig. 10 — Cellulose acetate fiber showing extensive surface erosion after one year of marine exposure. Original magnification 1000X. Photo by F. G. Foster.

hand side of the rod. In this case, the pholads obtained their start in the jute wrapping and then were able to progress into the Lucite. There was no attack evident on portions of the rod which were not wrapped with jute. Also, it is interesting to note that the jute in this particular case was treated with an impregnant consisting of 50 parts asphalt, 50 parts paraffin and 2 parts zinc naphthenate. Although this mixture was not highly effective as a preservative, it did serve to hold the jute in position long enough to enable the borers to become established. There has been no evidence of penetration of Lucite rods on which untreated jute was wrapped. Here, apparently, the jute was destroyed by microbiological attack or other borers such as limnoria before the pholads could become well established.

In this progress report no detailed comparison of the performance of natural fibers, such as jute treated with various preservatives, will be attempted. Results in many cases are still inconclusive. As might be anticipated, however, the jute specimens as a group have suffered much heavier deterioration by borers and microorganisms than the plastics, elastomers and casting resins. Particularly noteworthy is the fact that although there is considerable evidence of bacterial attack upon micro-

scopic examination, there is also much degradation evident by the fungi. Pin holes in the cell walls with associated fungal hyphae are extensive.

Secondary cellulose acetate has been quite susceptible to microbiological deterioration. Yarn has been destroyed in just six months of marine exposure, not by borers, but predominantly by bacteria. Upon microscopic examination, the fibers show severe surface erosion due apparently to bacterial attack. In the marine samples which have been stained and examined, hyphae have been evident in only one isolated instance. The extent of the pitting and surface erosion in one of the marine samples after one year in test can be seen in Fig. 10. This characteristic pattern of erosion is also evident in samples of cellulose acetate yarn from soil burial. Such a sample after 60 days of burial is shown in Fig. 11. Here again attack appears to be predominantly of bacterial origin.

In the marine borer tests, the sample rods become heavily fouled in the water area, as may be noted in Fig. 12. Although these fouling organisms do not use the materials to which they attach as a source of food, it is well known that they can do mechanical damage or chemically influence the environment beneath them. The reader who is particularly interested in the broad subject of fouling and its effect on materials,

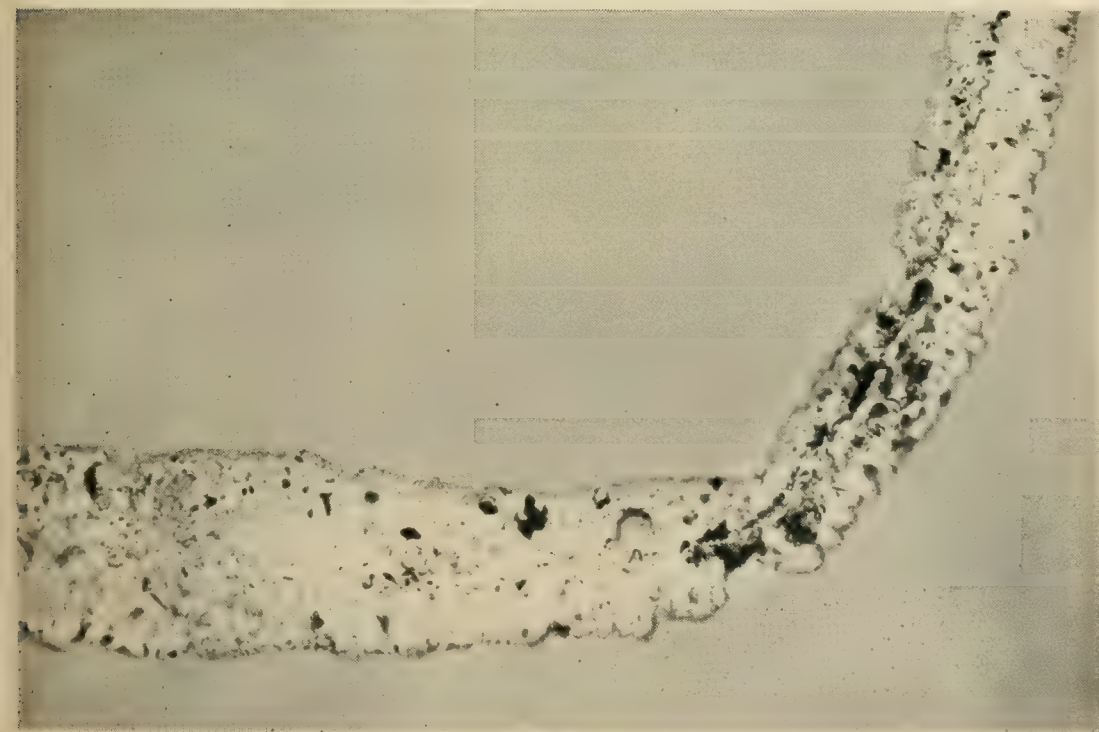


Fig. 11 — Cellulose acetate after 60 days in laboratory soil burial. Note characteristic surface erosion comparable to that shown in Fig. 10 for marine test sample. Original magnification 500X. Photo by F. G. Foster.



Fig. 12 — Test rack being lifted from water at Wrightsville Beach. Heavy accumulation of fouling on test rods in water-exposed area stops at point where rods entered the sediment.

structures and coatings, is referred to the report of the investigations conducted at the Woods Hole Oceanographic Institution during the years 1940 to 1946.¹⁷ The restricted areas beneath fouling, particularly under the bases of calcareous organisms such as barnacles, provide ideal cells for bacterial activity. Conditions of pH and aeration may be markedly different in these confined areas from those in the surrounding water. Some of the test rods made of polyvinyl chloride plastics containing basic lead stabilizers, illustrate the fact rather dramatically. Anaerobic, sulfate-reducing bacteria are common marine organisms which release hydrogen sulfide in the process of breaking down organic material. Under tightly adhering fouling, aerobic bacteria can utilize the free oxygen much more rapidly than it can be replaced by diffusion from the surrounding water. Once the oxygen has been depleted, the anaerobic organisms begin their activity and cause relatively high concentrations of hydrogen sulfide to be built up. The hydrogen sulfide reacts with the basic lead salts used as stabilizers and produces black lead sulfide. The sharp boundaries of the different environmental conditions existing beneath the base of a barnacle on one of the polyvinyl chloride test rods are illustrated in Fig. 13. Here the pattern of the barnacle base has literally been reproduced by the sulfiding which occurred under it. The black border and black radiating lines correspond to areas of exception-

ally close contact. The radial extent of this sulfiding in the bottom end of the rod which was embedded in the sediment, as compared to the top or water end, is shown in Fig. 14. It must be emphasized that there has been no indication to date of any adverse effect on the physical properties of plastics which have been sulfided in this way.

VI. CABLE SAMPLES FROM SERVICE

The samples of submarine cables which have been examined to the present time represent both telegraph and telephone cables. The samples of telegraph cable have been obtained through the cooperation of the Western Union Telegraph Company. It takes considerable time to assemble a large number of specimens during the course of routine repair operations. As a result, although some 22 different samples, the majority from the North Atlantic, have been examined, it is possible to make only

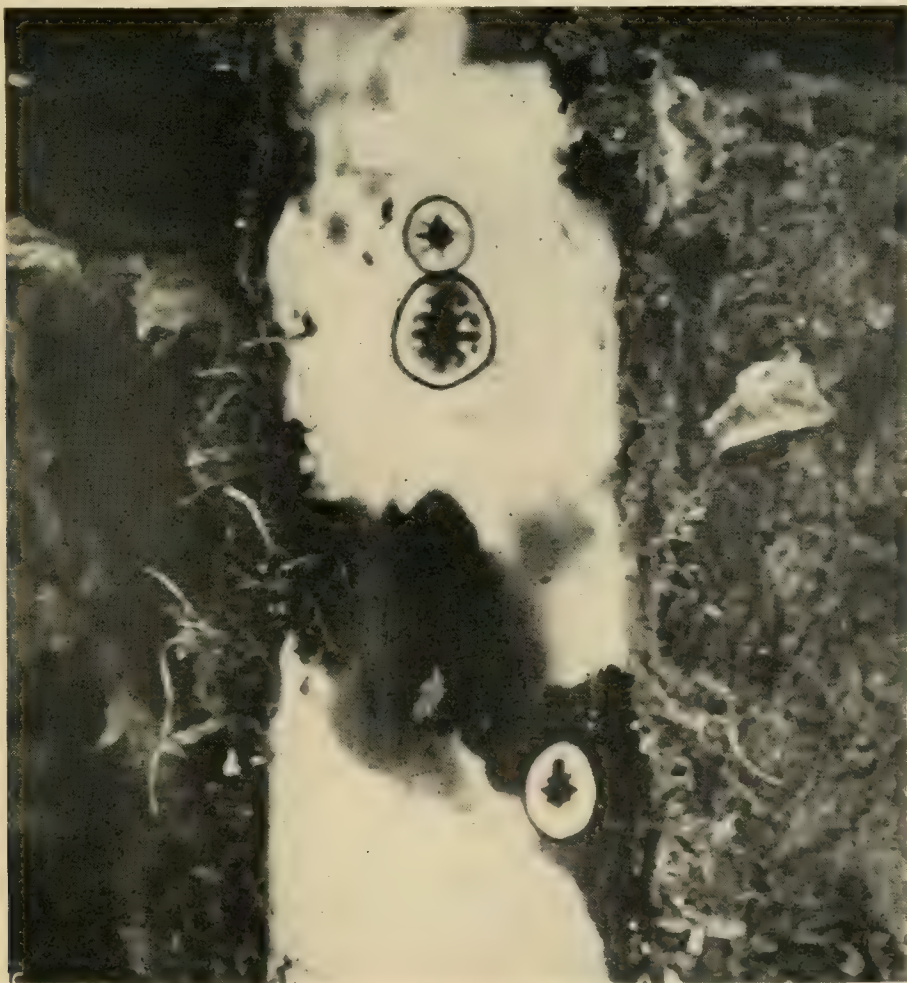


Fig. 13 — Sulfiding of Poly(Vinyl Chloride) plastic test rod beneath barnacles. The black, circular border and center area represent sulfiding at points of exceptionally close contact. Original magnification 2X. Photo by J. B. DeCoste.

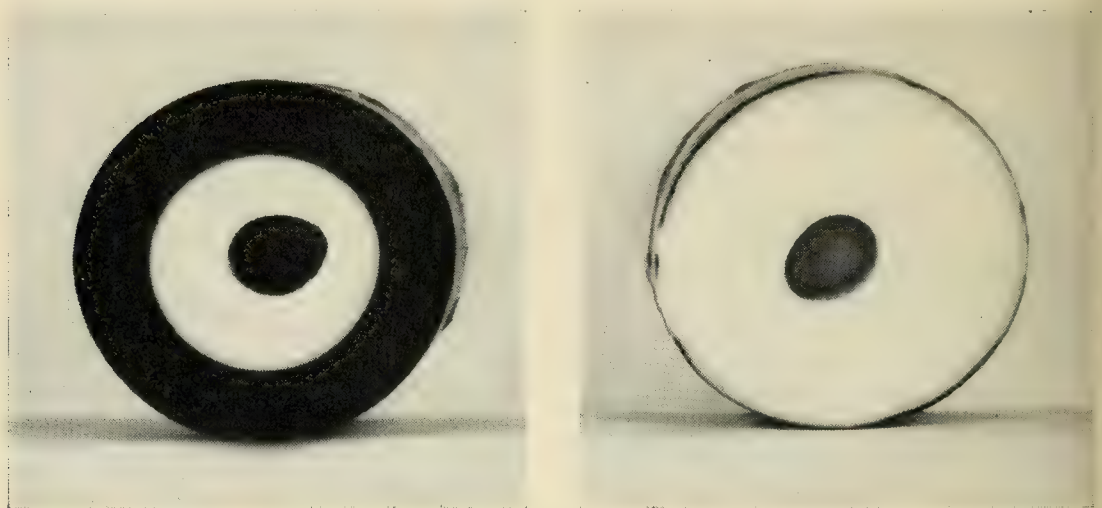


Fig. 14 — Cross-section of Poly(Vinyl Chloride) rod showing sulfiding due to sulfate-reducing marine bacteria. Bottom end was in sediment — top in water. Exposed at Daytona Beach for two years. Original magnification 1.6X.

general observations and broad comparisons. The locations and depths from which the samples were obtained are given in Table XIV. In most cases, a sample 3-feet long is obtained. Twenty-two different pieces of such size represent a rather small sample with respect to the total marine environment. Here again the program assumes more value as additional samples are obtained.

The procedure employed in examining one of these cable sections is about as follows. First, the over-all external condition is observed and recorded. Then, the outer jute covering, armor wires, inner jute bedding, and cloth and metallic tapes, if any, are removed progressively. The armor wires are examined in detail by electrochemists to determine the extent and kind of corrosion. The jute is tested in various ways. If sufficient material is available tensile strength is measured. Microscopic examination of representative fibers is also made. In the case of the inner jute which is not treated with tarry materials, damage counts are run according to the procedure of MacMillan and Basu.¹⁸ According to this method, deoiled and dewaxed fibers are permitted to swell in 10 per cent sodium hydroxide solution. Following this they are treated in a bath of 133 per cent weight to volume, aqueous zinc chloride solution over steam. Undamaged fibers swell as tight helices while damaged fibers swell as bundles of parallel fibrils. The fibers are mounted in zinc chloride solution on a microscope slide and counted to determine the per cent of damaged fibers.

To examine the integrity of the insulation on the conductor, the electrolytic procedure of Blake, Kitchin and Pratt¹⁹ is used. An electrolytic

cell is set up with 20 per cent copper sulfate solution as the electrolyte, and a copper plate as the anode. The cathode is a loop of the insulated conductor from the cable. Any plating out of copper on the cathode indicates a break in the insulation. In this way the integrity of a relatively long length of conductor can be examined critically and simply.

One of the outstanding facts apparent from the examination of cable samples has been the evident importance of the outer jute in limiting the corrosion of armor wires. Galvanized steel armor wires which still retain the protection of flooding compounds, such as asphalt, tar or pitch, together with outer servings of impregnated jute, have shown negligible steel corrosion within 40 years, and in one case for as long as 66 years. On the other hand, most of the corrosion of armor wires which has been observed has occurred in cable from which a major part, or all, of the outer jute has been lost.

TABLE XIV — LOCATIONS AND DEPTHS FROM WHICH SUBMARINE
CABLE SAMPLES HAVE BEEN OBTAINED FOR
LABORATORY EXAMINATION

BTL No.	Location		Depth
	<i>Latitude</i>	<i>Longitude</i>	<i>Fathoms</i>
110	approx. 81° 30' 00"N	24° 30' 00"W	5 approx.
111	approx. 81° 30' 00"N	24° 30' 00"W	5 approx.
112	approx. 81° 30' 00"N	24° 30' 00"W	5 approx.
113	approx. 81° 30' 00"N	24° 30' 00"W	5 approx.
114	approx. 81° 30' 00"N	24° 30' 00"W	5 approx.
135a	23° 45' 00"N	81° 57' 30"W	830
136	Several miles south of Long Island, N. Y. Exact location unknown.		90
137	40° 13' 40"N	71° 07' 25"W	90
163	48° 36' 06"N	36° 23' 36"W	2460
164	51° 53' 42"N	10° 37' 18"W	54
165	51° 40' 42"N	13° 02' 12"W	630
166	51° 55' 21"N	11° 58' 18"W	387
167	47° 52' 24"N	38° 23' 12"W	2460
168	47° 22' 06"N	42° 14' 12"W	2175
169	36° 41' 46"N	25° 38' 09"W	1180
195	45° 28' 21"N	60° 20' 24"W	106
196	46° 41' 36"N	56° 18' 18"W	37
197	47° 00' 32"N	56° 51' 40"W	100
200	44° 25' 40"N	63° 25' 15"W	46
212	45° 08' 38"N	54° 33' 06"W	82
281	43° 38' 10"N	55° 07' 00"W	2090
282	39° 17' 24"N	70° 12' 15"W	1450
283	53° 57' 00"N	165° 50' 00"W	Unknown

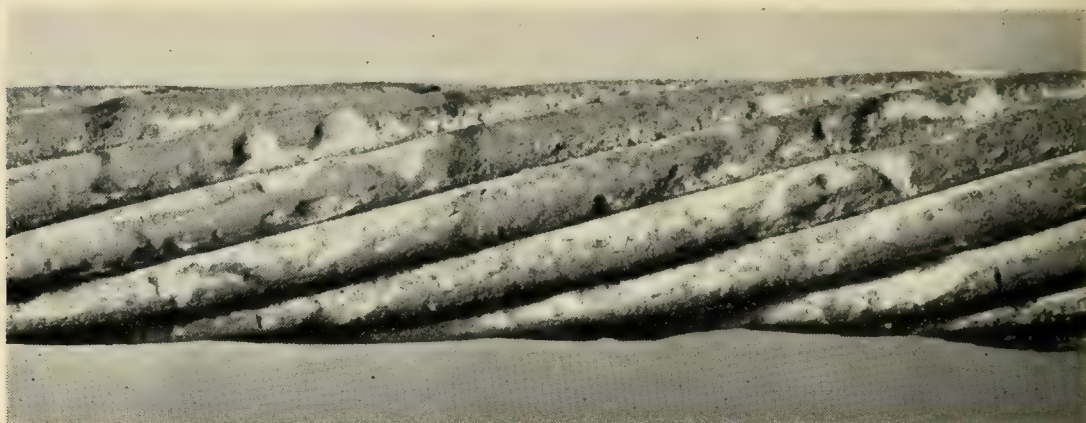


Fig. 15 — Corrosion pockets in galvanized steel armor wires of submarine cable after 12 years of service.

Of special interest is a particular type of corrosion which has occurred in two separate cable samples — one from Alaskan waters, the other from off the southern coast of Newfoundland. In one case the age of the cable was 12 years, and in the other 36 years. The Alaskan sample was located in an area characterized by water velocities of 5 to 9 knots. In both instances, the outer jute and most of the flooding compound was gone. Corrosion, instead of starting and progressing on the outer surface of the wires, had started, and been confined largely to the sides of the wires. Usually there are corresponding areas of corrosion on two adjacent wires to form "corrosion pockets." These pockets are illustrated in Fig. 15. In the case of the 12 year old cable, corrosion caused failure of the armor wires. In the case of the cable which was in service 36 years, failure was reported to have occurred from chaffing on a rocky bottom. Close examination suggests that failure may more reasonably be attributed to severe corrosion of the type just described. The exact cause of the corrosion pattern is still to be determined.

Sufficient samples have not been examined as yet to form a coordinated picture with respect to the inner jute. In the case of cable samples from the North Atlantic, the inner jute bedding was in good condition in cables which had been in service for as long as 30 or 40 years. Samples more than 40 years old showed the effect of deterioration. Although only a limited number of samples have been examined from the Caribbean, most of them from water about 50 feet deep, jute and cotton tape components were in poor condition in certain spots. It was evident that microbiological deterioration of the jute had occurred. In no case has there been any evidence of deterioration of the insulation of the central conductor of any of the cable samples. In the case of the older cable samples the insulation was gutta percha, but in the most recent samples it has been polyethylene.

VII. SUMMARY

A progress report has been presented on the results of a test program designed to determine the relative resistance of materials to marine biological attack. Specific test results have been reported wherever possible, predominantly from the laboratory test procedures. In the case of the natural exposure tests, which are intended to provide correlative data for the laboratory program over longer periods of time, the information which has been assembled thus far is of a more general nature. There follows a summary of the more important information which has been obtained:

1. In the biochemical oxygen demand-type test it has been found that polyethylene is not utilized by the aerobic bacteria or the anaerobic sulfate-reducing bacteria. Polyvinyl chloride plastics are attacked according to the way in which they are plasticized. All of the samples tested which had an added external plasticizer, including the rigid plastic, were attacked to some degree. In the latter case the attack was apparently due to lubricants. The semi-flexible polyvinyl chloride copolymers, and the polyvinyl chloride resin alone, were not utilized by the bacteria. The five elastomers assayed were all attacked by aerobic bacteria, neoprene being the most resistant. The epoxide casting resin did not serve as a source of carbon for the organisms, but further testing is required with a polyester casting resin.

2. Coiled conductors insulated in one case with a rigid polyvinyl chloride, and in the other with GR-S, have been exposed half in sea water and half in marine sediment in the laboratory for thirteen months. Capacitance measurements show that a considerable change has occurred in the GR-S insulation apparently as a result of bacterial attack. Although there has been a slight rise in the capacitance values for the polyvinyl chloride-insulated conductors during the last five months, further observations are necessary before attack can be considered definite.

3. In three years of actual marine exposure of plastics, elastomers and casting resins, there have been definite penetrations by marine borers of only three materials — a test rod of silicone rubber, a 0.0035-inch film of polymonochlor-trifluoroethylene wrapped on a Lucite rod, and on Lucite rods themselves. The first two cases represent single instances of penetration — both by pholads. The Lucite rods were penetrated at many places by pholads as a result of the organisms getting started in an asphalt-impregnated jute wrapping and then progressing into the Lucite.

Secondary cellulose acetate yarn and tow have been deteriorated badly, apparently by bacteria, in as short a time as six months.

Natural fibers, notably jute, have been degraded extensively by borers and microorganisms. There is considerable evidence of fungus attack.

Under fouling and in the sediment area, rods of polyvinyl chloride plastics containing basic lead stabilizers have been blackened as the result of hydrogen sulfide produced by sulfate-reducing bacteria reacting with the lead salts to give black lead sulfide. This sulfiding has caused no apparent degradation of the physical properties of the plastics.

4. The examination of cable samples from service has indicated that the impregnated outer jute serves an important function in limiting corrosion of armor wires. Generally, when corrosion is present the outer jute has been lost.

Two unusual cases of extensive corrosion have been reported — one in a cable 12 years old, the other in a cable which was in service 36 years. In both cases, corrosion occurred in pockets between adjacent armor wires rather than on the outside surfaces (water side) of the wires.

The performance of the inner jute in samples from service has been generally good for as long as 30 or 40 years in deep water. In samples from relatively shallow water in the Caribbean, inner jute bedding was badly deteriorated in as short a time as five years.

ACKNOWLEDGMENTS

The data which have been acquired in this program are the result of teamwork by many different members of the Laboratories. Special thanks are due to Priscilla Leach for obtaining much of the data in the BOD test, and her general assistance on many phases of the program. Madeline L. Cook is responsible for the majority of the electrical measurements in the conductor tests. T. D. Kegelman, formerly of the Laboratories, and W. C. Gibson gave considerable assistance in setting up the equipment and establishing the procedures for measurement of capacitance and conductance. Many members of the Chemical Research Department cooperated in furnishing the materials which have been tested. J. B. DeCoste has been particularly helpful in assembling the required information on the materials. From outside the Laboratories, Professor Claude E. ZoBell of the Scripps Institute of Oceanography has made many helpful suggestions and comments with regard to the laboratory portion of the program, and provided supplementary enrichment cultures and samples of sediment. The marine borer test program has been executed with the cooperation of A. P. Richards, William F. Clapp Laboratories, Inc., who has furnished much valuable information in many discussions of the tests. The cooperation of C. S. Lawton, Western Union Telegraph Company, and the personnel of C. S. Lord Kelvin in

obtaining samples of telegraph cable from service is gratefully acknowledged.

REFERENCES

1. C. Chilton, The Gribble (*Limnoria lignorum*, Rathke) Attacking a Submarine Cable in New Zealand. *Annals and Mag. Natural History*, Ser. 8, **18**, p. 208, 1916.
2. G. E. Preece, On Ocean Cable Borers, *Telegr. Jour. and Electr. Rev.*, **3**, pp. 296-297, 1875.
3. E. Jona, I cavi sottomarini dall'Italia alla Libia, *Atti della Soc. ital. per il Progr. delle Sci.*, **6**, pp. 263-292, 1913.
4. R. J. Menzies and Ruth Turner, The Distribution and Importance of Marine Wood Borers in the United States, presented at Second Pacific Area National Meeting, A. S. T. M., Paper **93**, 1956.
5. W. F. Clapp, and R. Kenk, *Marine Borers, A Preliminary Bibliography*, Parts I and II, The Library of Congress, Tech. Inform. Div., Washington, D. C., 1956.
6. F. Roch, Die *Teredinidae* de Mittelmeeres. *Thalassia* **4**, pp. 1-147, 1940.
7. Notes on Everyday Cable Problems. Distribution of Electricity, **8**, pp. 1896-1898, W. T. Henley's Telegraph Works Co., Ltd., London, November, 1935.
8. P. Bartsch and H. A. Rehder, *The West Atlantic Boring Mollusks of the Genus Martesia*, Smithsonian Inst. Misc. Collections 104, Washington, D. C., **11**, 1945.
9. L. R. Snoke and A. P. Richards, Marine Borer Attack on Lead Cable Sheath, *Science*, **124**, p. 443, 1956.
10. C. E. ZoBell, *Marine Microbiology*, Chronica Botanica Co., Waltham, Mass., 1946.
11. R. Y. Morita, and C. E. ZoBell, Bacteria in Marine Sediments. Off. of Naval Res., Research Reviews, p. 21, July, 1956.
12. C. E. ZoBell and R. Y. Morita, Effects of High Hydrostatic Pressure on Physiological Activities of Marine Microorganisms, Off. of Naval Res., Contr. N6 onr-275 (18) Project NR 135-020, Semi. Ann. Prog. Rept., July, 1955.
13. C. E. ZoBell and Josephine Beckwith, The Deterioration of Rubber Products by Micro-Organisms, *Jour. Amer. Water Works Assoc.*, **36**, pp. 439-453, 1944.
14. C. E. ZoBell, Action of Microorganisms on Hydrocarbons, *Bact. Rev.*, **10**, Nos. 1-2, March-June, 1946.
15. E. S. Barghoorn and D. H. Linder, Marine Fungi, Their Taxonomy and Biology, *Farlowia*, **1**, pp. 395-467, Jan., 1944.
16. S. P. Myers, Marine Fungi in Biscayne Bay, Florida. *Bull. of Marine Sci. of the Gulf and Caribbean*, **2**, pp. 590-601, 1953.
17. *Marine Fouling and Its Prevention*. United States Naval Institute, Annapolis, Md., 1952.
18. W. G. MacMillan and S. N. Basu, Detection and Estimation of Damage in Jute Fibers — Part I: A New Microscopic Test and Implications of Certain Chemical Tests. *Jour. Text. Inst.*, **38**, pp. T350-T369, 1947.
19. J. T. Blake, D. W. Kitchin and O. S. Pratt, The Microbiological Deterioration of Rubber Insulation. Presented at A.I.E.E. General Meeting, New York, Jan., 1953.
20. P. Deschamps, Xylophaga Marins. Protection de Bois Immergés contres les Animaux Perforants. *Peint.-Pigm.-Vern.*, **28**, pp. 607-610, 1952.
21. C. E. ZoBell, Some Effects of High Hydrostatic Pressure on Physiological Activities of Bacteria, *Proc. Soc. Amer. Bact.*, pp. 26-27, 1955.

Dynamics and Kinematics of the Laying and Recovery of Submarine Cable

By E. E. ZAJAC

(Manuscript received June 5, 1957)

This paper is an attempt to formulate a comprehensive theory with which the forces and motions of a submarine cable can be determined in typical laying and recovery situations. In addition to the fundamental case of a cable being laid or recovered with a ship sailing on a perfectly calm sea over a horizontal bottom, the effects of ship motion, varying bottom depth, ocean cross currents, and the problem of cable laying control are considered. Most of the results reduce to simple formulas and graphs. Their application is illustrated by examples.

TABLE OF CONTENTS

	Page
I. Introduction.....	1132
II. Basic Assumptions.....	1133
III. Two-Dimensional Stationary Model.....	1134
3.1 General.....	1134
3.2 Normal Drag Force and the Cable Angle α	1135
3.3 Tangential Drag Force.....	1139
3.4 Sinking Velocities and Their Relationship to Drag Forces.....	1141
3.5 General Solution of the Stationary Two-Dimensional Model.....	1143
3.6 Approximate Solution for Cable Laying.....	1146
3.7 Approximate Solution for Cable Recovery.....	1149
3.8 Shea's Alternative Recovery Procedure.....	1153
IV. Effects of Ship Motions.....	1154
4.1 Tensions Caused by Ship Motions.....	1154
V. Deviations from a Horizontal Bottom.....	1158
5.1 Kinematics of Laying Over a Bottom of Varying Depth.....	1158
5.2 Time-Wise Variation of the Mean Tension in Laying Over a Bottom of Varying Depth.....	1161
5.3 Residual Suspensions.....	1163
VI. Cable Laying Control.....	1165
6.1 General.....	1165
6.2 Accuracy of the Piano Wire Technique.....	1166
VII. Three-Dimensional Stationary Model.....	1169
7.1 General.....	1169
7.2 Perturbation Solution for a Uniform Cross-Current.....	1172
Appendix A. Discussion of the Two-Dimensional Stationary Configuration for Zero Bottom Tension.....	1175
Appendix B. Computation of the Transverse Drag Coefficient and the Hydrodynamic Constant of a Smooth Cable from Published Data.....	1177

Appendix C. Some Approximate Solutions for Laying and Recovery.....	1180
C.1 Laying.....	1180
C.2 Recovery.....	1183
Appendix D. Analysis of the Effect of Ship Motion.....	1184
D.1 Formulation of the Differential Equations.....	1184
D.2 Perturbation Equations.....	1187
D.3 Solution of the Perturbation Equations.....	1189
D.4 Transverse Response.....	1190
D.5 Second-Order Longitudinal Response.....	1191
D.6 Numerical Results.....	1194
Appendix E. Tension Rise with Time for Suspended Cable.....	1195
E.1 Formulation of the Solution of the Problem.....	1195
E.2 Nomograph for the Solution of Equation (99).....	1198
E.3 Numerical Example.....	1199
Appendix F. The Three-Dimensional Stationary Model.....	1202
F.1 Derivation of the Differential Equations.....	1202
F.2 Perturbation Solution for a Uniform Cross Current.....	1204
Acknowledgments.....	1206
References.....	1206

GLOSSARY OF SYMBOLS

A	Amplitude of harmonic ship motion
$c_1, c_2, \bar{c}_2$	Longitudinal wave velocity; transverse wave velocity in air and water
C_D, C_f	Transverse and tangential drag coefficients
d	Cable diameter, also distance behind the ship at which the cable enters the lower stratum
D_N, D_T	Normal and tangential unit drag forces
e	Sidewise distance from the laid cable to the ship
$\overline{EA}$	Extensile rigidity
h	Ocean depth
$\bar{h} = \frac{wh}{\overline{EA}}$	Dimensionless ocean depth
H	Hydrodynamic constant
L	Inclined cable length from surface to bottom, also from ship to surface
N_R	Reynolds number
p, q	Longitudinal and transverse deviational cable displacements
P_0, Q_0	Longitudinal and transverse ship displacements

P_1	Deviation from mean pay-out or haul-in rate
S, X	Arc length and horizontal distance from the touchdown point to the ship
$\bar{S} = \frac{S}{h}, \quad \bar{X} = \frac{X}{h}$	Dimensionless forms of S and X
t	Time
$\bar{t} = \frac{Vt}{h}$	Dimensionless time
T, T_s, T_0	Cable tension at an arbitrary point, at the ship, and at the bottom
$\bar{T} = \frac{T}{wh}, \quad \bar{T}_s = \frac{T_s}{wh}, \quad \bar{T}_0 = \frac{T_0}{wh}$	Dimensionless forms of T, T_s and T_0
T_p, T_q	Cable tension due to longitudinal and transverse ship motion
V, V_c	Ship speed, pay-out or haul-in rate
V_N, V_T	Normal and tangential velocity of the water relative to the cable configuration
V_t	Tangential velocity of the water relative to a cable element
w, w_a	Submerged and in-air unit cable weight
α, α_0	Critical angle, approximate critical angle
α_s	Cable angle at the surface
β	Descent angle, cross current orientation (Section 7.1)
$\gamma = \frac{2 - \sin^2 \alpha}{\sin^2 \alpha}$	Constant, also ascent angle
ϵ	Slack
θ	Orientation of a cable element
θ, ψ	Spherical polar coordinates for the three-dimensional model
κ, λ	Constants (see Appendix C)

$\Lambda = \frac{C_D \rho d V^2}{2} = \frac{\cos \alpha}{\sin^2 \alpha}$	Constant
μ, ν	Constants
ν	Kinematic viscosity
ξ, η, ζ	Rectangular coordinates for the three-dimensional model
ρ	Mass density of water
ρ_c, ρ_w	Mass per unit length of cable in air and water
ϕ	Deviation from the stationary angle, also angle between ξ axis and $\vec{V}$ (Section 7.1)

I. INTRODUCTION

In the summer of 1857, the first attempted laying of a transatlantic cable ended dismally when, after only a few hundred miles had been laid, the cable broke and fell into the sea. Although fouling of the payout gear caused by a negligent workman was the principal suspected reasons for the failure, its occurrence aroused great interest in the detailed dynamics and kinematics of the laying of submarine cable, and leading British scientists such as Kelvin and Airy published analyses of this problem in late 1857 and early 1858.^{1, 2, 3, 4, 5}

However, after this initial activity, interest in submarine cable dynamics and kinematics evidently waned for there appear only sporadic subsequent investigations in the literature.^{6, 7, 8, 9, 10} Further, the results of the early and subsequent analytical investigations have been, by and large, little utilized in cable laying and recovery practice. One can conjecture several reasons for this. For one, because the early analytical work was done before the advent of modern hydrodynamic theory, it did not rest on a secure base. Thus, as late as 1875, one finds vigorous debate over the nature of the tangential resistance of water to the cable.⁹ For another, the results of the analyses could not all be expressed in terms of elementary functions and required the numerical evaluation of some definite integrals. In the 1850's this was a tedious and laborious process. However, these are probably secondary reasons. For, after another failure in the early summer of 1858, a transatlantic cable was successfully laid in August of that year. The mechanical problem of depositing a cable was thus proved surmountable without complicated mathematical analyses, and the marriage of analysis and practice was never fully realized.

However, a present-day submerged-repeater transoceanic cable is a delicate and expensive transmission system. Reducing the amount of cable deposited by as little as one per cent can result in a substantial saving in the first cost of such a system. Its repair is a costly operation requiring the sustenance of an ocean ship and its crew. Therefore, it is important to lay the cable without wasteful excess and with minimum chances for failure after laying. Further, it is important that repair, if necessary, be as efficient as possible. To accomplish these things, an understanding of the dynamics and kinematics of cable laying and recovery is essential.

The purpose of this paper is to provide some of this understanding in as straightforward a way as possible. To this end concepts and results are stressed in the main part of the paper, mathematical details being given in the appendices. Moreover, we hope to show that the results of the analysis can provide a numerical basis for decision making in many of the laying and recovery operations. Most of these results can be expressed in the form of simple formulas and graphs. Several numerical examples are included to illustrate concretely how the results can be applied in practice.

The general plan of the paper is to proceed from simple to more refined models of the laying and recovery processes. Thus, we discuss first what we have called the *two-dimensional stationary model*. This model is appropriate for laying and recovery on or from a perfectly flat bottom while sailing on a perfectly still sea. As a preliminary to this discussion, we consider in some detail the hydrodynamic behavior of typical deep sea submarine cable. We then take up the effects of the ship motions which are induced by wave action and the effects of a bottom of varying depth. These considerations are followed by a short discussion of the problem of controlling the cable pay-out properly during laying and the associated problem of the accuracy of the present taut wire method of determining ship speed. Finally, we consider the *three-dimensional stationary model* and the effects of ocean cross currents.

II. BASIC ASSUMPTIONS

Our analyses, like most analyses of physical problems, are based on idealizations or mathematical models of the actual physical system. The extent of validity of these models must be ultimately determined by experiment and experience. However we shall try to give the reader an idea of when they are clearly applicable and when they are not.

All of the models we consider contain two basic idealizations, namely,

- (1) No bending stiffness in cable, i.e., it is a perfectly flexible string,
- (2) The average forward speed of the ship is constant.

Bending effects are caused by locally large curvatures, and are significant mainly where the cable leaves the pay-out sheaves and at the ocean bottom. However, for a cable with a steel strength member, bending even to the small radius of the pay-out sheave typically does not materially reduce the tension required to break the cable. Hence, in these cases we can expect an analysis based on the first idealization to give a reasonable idea of when cable rupture will occur. In laying, ship speeds are normally steady and, with the exception of the fluctuations caused by wave action which we consider later in the paper, the second idealization is reasonable also. In recovery, on the other hand, ship speeds are apt not to be steady, and the second idealization is more tenuous. But because of the very slow speeds usually employed, this idealization may in fact be meaningful in recovery as well.

III. TWO-DIMENSIONAL STATIONARY MODEL

3.1 General

Assume that the cable ship is sailing at a constant horizontal velocity, that the cable pay-out or haul-in rate is constant, and that the drag of the water on the cable depends only on the relative velocity between the water and the cable. Further, assume that in a frame of reference translating with the ship the cable configuration is time-independent or stationary. This idealized model of the cable laying or recovery process we call the *two-dimensional stationary model*.

This is the model which has been considered in the previous analytical studies.¹⁻¹⁰ As the early investigators quickly pointed out, when the tension at the bottom of the cable is zero, the cable, according to this model, can lie in a straight line from ship to ocean bottom. During laying, when slack is normally paid out, the zero tension condition actually occurs, and hence this case is of considerable practical importance.

The straight line can in fact be shown to be the *only* solution which can satisfy all the observed boundary conditions. This point is discussed in detail in Appendix A. That the straight line is a possible configuration can be seen from Fig. 1. In the vector diagram the velocity of the water with respect to the cable is resolved into a component V_N normal to the cable and a component V_T tangential to it. Associated with V_N and V_T are normal and tangential water resistance or drag forces D_N and D_T . In the straight line configuration, the cable inclination is such that D_N just balances the normal component of the cable weight forces. The situation is thus analogous to that of a chain sliding on an inclined plane, with the forces D_N corresponding to the normal reaction forces of the plane. Summing forces in the normal direction, we get, therefore,

$$w \cos \alpha = D_N, \quad (1)$$

while the summation in the tangential direction gives for T_s , the tension at the ship,

$$T_s = wL \sin \alpha - D_T L. \quad (2)$$

Here w is the weight per unit length of immersed cable, α is the cable's angle of incidence, D_N and D_T are the normal and tangential drag forces per unit length respectively, and L is the inclined length of the cable. For most submarine cable used currently the force $D_T L$ is negligible and we arrive at

$$T_s \approx wL \sin \alpha = wh, \quad (3)$$

where h is the ocean depth at the cable touchdown point. Hence, during slack laying the cable tension at the ship is very nearly equal to the weight in water of a length of cable equal to the ocean depth.

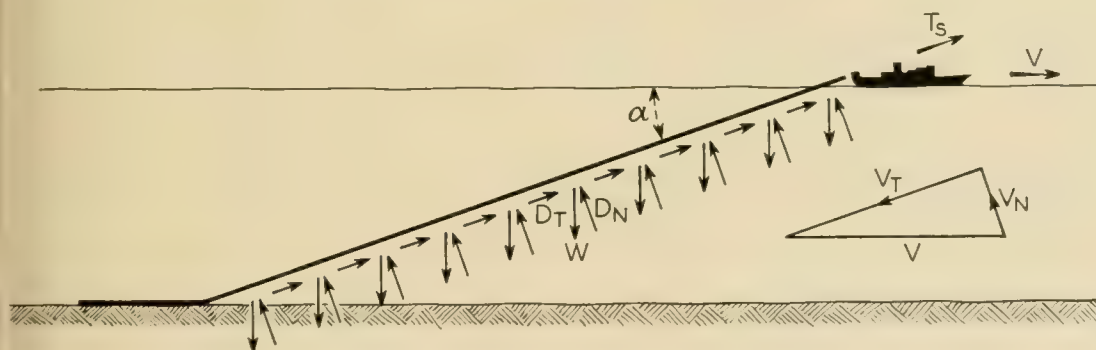


Fig. 1 — Forces acting on a cable in normal laying.

The straight-line solution is the simplest and probably the most important result to be obtained from the stationary two-dimensional model. We shall derive results for other important situations from this model also. As a preliminary, we study first, however, the nature of the normal and tangential drag forces D_N and D_T .

3.2 Normal Drag Force and the Cable Angle α

The resistance at sufficiently slow speeds to the flow of a fluid around an immersed body varies as the square of the fluid velocity. This relationship is usually written as*

$$D_N = C_D \frac{\rho V_N^2 d}{2}, \quad (4)$$

* For towed stranded wire experimental verification of this relationship is reported in Reference 11.

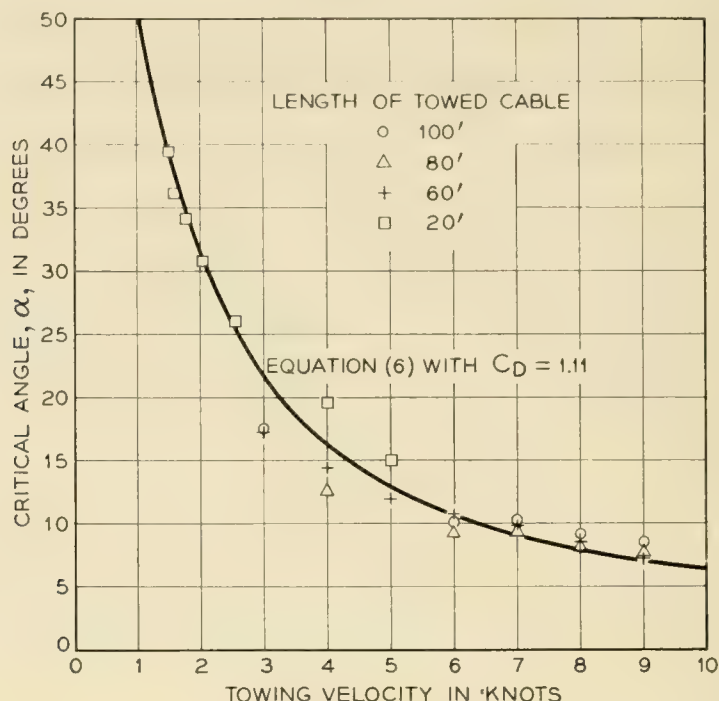


Fig. 2 — Experimental and theoretical variation of critical angle with towing velocity for cable No. 1.

where D_N is the normal drag force per unit length, C_D is the so-called drag coefficient, ρ is the mass density of the fluid, and d is the diameter of the cable. For the straight-line configuration, the vector diagram in Fig. 1 shows that

$$V_N = V \sin \alpha. \quad (5)$$

Substitution of (5) and (4) into (1) yields in turn

$$w \cos \alpha = \frac{C_D \rho V^2 d}{2} \sin^2 \alpha. \quad (6)$$

Equation (6) suggests how the value of the drag coefficient C_D can be obtained experimentally. By towing a length of cable in water at a constant velocity, one can establish the straight-line configuration. The angle α can then be measured as a function of velocity, from which C_D can be computed by (6).

Figs. 2 and 3 show the results of such tests together with plots of (6) for the indicated values of C_D . These results are taken from an analysis by A. G. Norem of experimental data obtained by H. N. Upthegrove, J. J. Gilbert, and P. A. Yeisley. The properties of these cables are listed in Table I.* To eliminate end effects different lengths of cable were towed,

* Cable No. 2 is very similar to present type D transatlantic telephone cable. For engineering calculations, type D can be considered the same as cable No. 2.

TABLE I — PROPERTIES OF CABLES No. 1 AND No. 2

Cable	No. 1	No. 2
Diameter (inches).....	0.75	1.25
Wt. in water (lbs/ft.).....	0.243	0.705
Outer covering.....	Polyethylene	Tar impregnated jute
Surface condition.....	Smooth	Rough
$\overline{EA}$ (twist restrained).....	—	4×10^6 lbs
$\overline{EA}$ (twist unrestrained).....	—	1.2×10^6 lbs

as is indicated by the plotted experimental points. It is seen that (6) gives a good fit to the experimental data over the entire velocity range.

If the cable has a smooth exterior, an estimate of the drag coefficient C_D can be computed from published values of resistance to flow about an immersed cylinder. This computation is described in Appendix B, where we have also tabulated computed values of C_D . For the smooth cable No. 1, the value of C_D obtained from Appendix B is 1.00 which is in fair agreement with the experimentally determined value of 1.11.

Although the drag coefficient C_D is a fundamental hydrodynamic parameter, it is not the most convenient description of the effect of the nor-

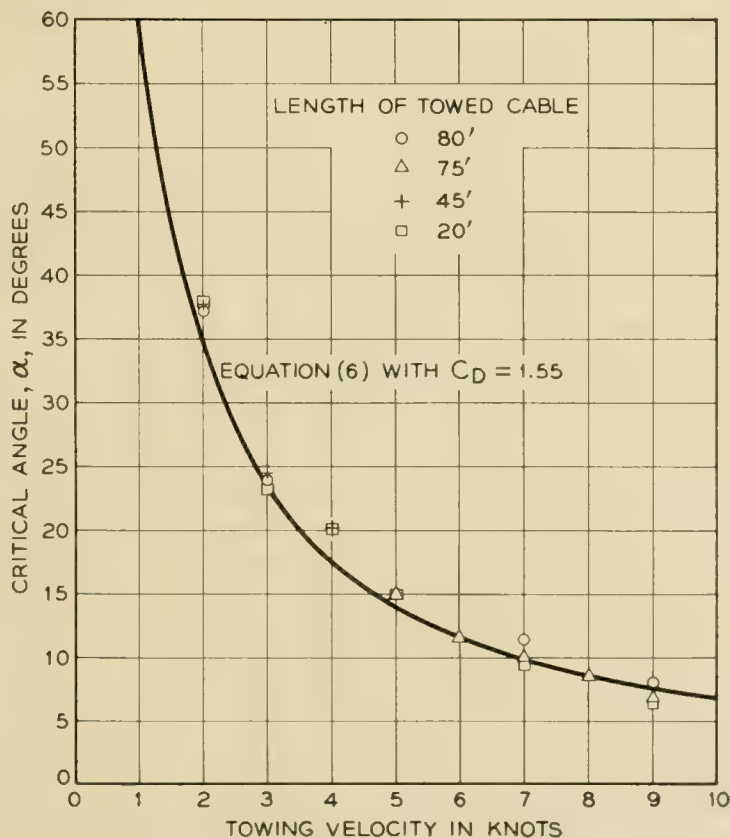


Fig. 3 — Experimental and theoretical variation of critical angle with towing velocity for cable No. 2.

mal component of water velocity. For small values of the incidence angle α

$$\cos \alpha \approx 1,$$

$$\sin \alpha \approx \alpha,$$

and (6) is approximately

$$\alpha_0 V = \left(\frac{2w}{C_D \rho d} \right)^{1/2}, \quad (7)$$

where α_0 is the approximate value of α . The quantity $(2w/C_D \rho d)^{1/2}$ is a constant for a given cable. It brings together all the cable parameters which influence the magnitude of the incidence angle α . If the angle α for a given speed is determined accurately, as can be done in a towing test or with a sextant during over-the-stern laying, this quantity is easily computed. Because of its importance, we shall call it the *hydrodynamic constant* of the cable and denote it by H , namely,

$$H = \left(\frac{2w}{C_D \rho d} \right)^{1/2}. \quad (8)$$

By virtue of (7) and (8) we may write

$$\alpha_0 V = H. \quad (9)$$

The constant H rather than the drag coefficient C_D will be used from this point on.

When the approximate relationship (9) is not valid, α can be obtained by solving (6). This gives

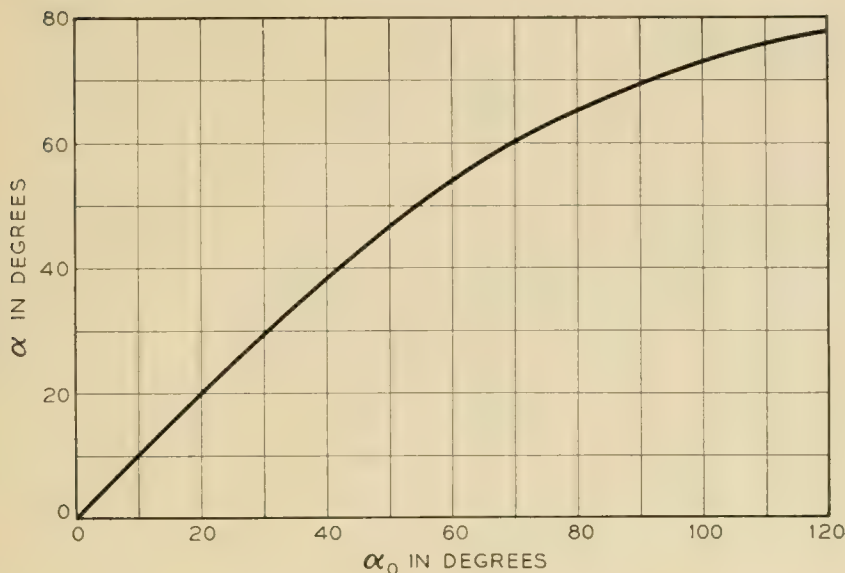
$$\cos \alpha = \sqrt{1 + \frac{1}{4} \left(\frac{H}{V} \right)^4} - \frac{1}{2} \left(\frac{H}{V} \right)^2, \quad (10)$$

where V is in knots and H in radian-knots. In terms of α_0 we obtain in turn

$$\cos \alpha = \sqrt{1 + \frac{1}{4} \alpha_0^4} - \frac{1}{2} \alpha_0^2. \quad (11)$$

This relationship is shown in Fig. 4, where the incidence angle α is plotted as a function of the approximate incidence angle α_0 . It is seen that for $\alpha < 20^\circ$ the difference between α_0 and α is negligible.

Physically α as given by (10) is the angle the cable assumes in the straight-line shape for the velocity V . However, in addition, (10) shows that α can be thought of as a dimensionless parameter which embodies

Fig. 4 — Variation of α with α_0 .

both the hydrodynamic properties of the cable and the ship speed. Thus, even when the configuration is not a straight line, we shall find it convenient to express results as a function of the single parameter α , rather than as a function of the two parameters H and V . For this reason, following Pode,^{11, 12} we call α the *critical angle*.

3.3 Tangential Drag Force

Over the range of velocities encountered in laying and recovery the drag coefficient C_D in equation (6) is essentially constant. However, the corresponding coefficient for the skin friction force associated with V_T , the component of flow along the cable, is not constant. For the cable of smooth exterior (cable No. 1), the expression

$$D_T = C_f \frac{1}{2} \rho V_t^2 \pi d, \quad (12)$$

with $C_f = 0.055/(N_R)^{0.14}$, was found to give good agreement with the experimental data, as is shown by Fig. 5. Here D_T is the skin friction or tangential drag force per unit length; V_t is the relative velocity of the water with respect to a cable element given for straight-line laying by

$$V_t = V_c - V \cos \alpha, \quad (13)$$

where V_c is the cable pay-out rate; ρ is the mass density of water; and N_R is the Reynolds number defined as $N_R = V_t L/\nu$, where ν is the kinematic viscosity of water. The data of Fig. 5 are for 100 foot lengths of cable towed in fresh water at a temperature of 60°F.

From (12) we find

$$D_T = \frac{0.055}{2} \rho \frac{\nu^{0.14} V_t^{1.86}}{L^{0.14}} \pi d. \quad (14)$$

This expression indicates that D_T for smooth cable depends on the inclined length of the cable as well as the relative tangential velocity V_t . The form of (13) suggests that the flow tangential to a smooth cable is similar to flow past a smooth plate. In such flow a turbulent boundary layer develops which grows in thickness with distance from the leading edge, resulting in a length dependence of the type shown by (14). Since Fig. 5 refers to 100 foot cable lengths, (13) is probably not accurate for the magnitudes of L occurring in deep-sea laying, and should be used only to obtain the order of magnitude of C_f .

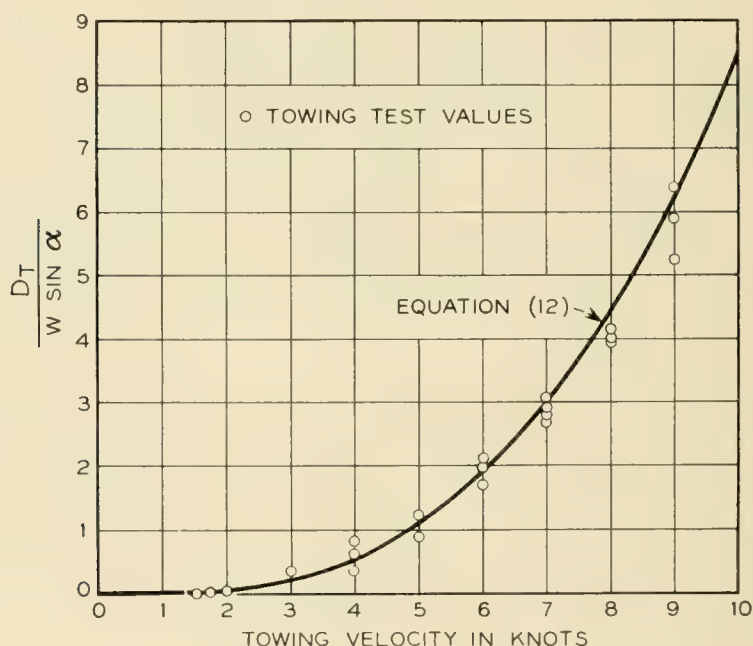


Fig. 5 — Experimental values of the tangential drag force for cable No. 1 compared with those obtained by equation (12).

For the cable with conventional jute outer covering, (cable No. 2), it was found that

$$D_T = 0.01 V_t^{1.48} \quad (15)$$

fits the experimental data obtained by towing test (Fig. 6). Whereas in (15) the constant 0.055 is dimensionless, the constant 0.01 in this equation has the dimensions necessary to give D_T in units of pounds per foot when V_t is in feet per second. We note that for this cable D_T is independent of the length of the cable.

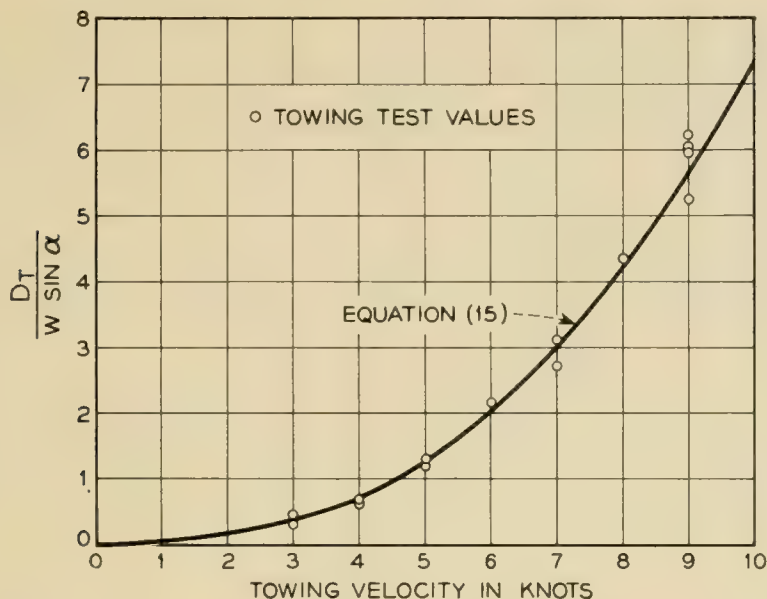


Fig. 6 — Experimental values of the tangential drag force for cable No. 2 compared with those obtained by equation (15).

The ratio of D_T to the tangential component of the cable weight force is given by $D_T/w \sin \alpha$. Equations (14) and (15) indicate that even for small values of α of the order of twelve degrees, $D_T/w \sin \alpha$ is of the order of 6 per cent for relative tangential velocities V_t of 1.0 feet/sec. In many situations V_t will be less than this value, and we can neglect D_T compared to $w \sin \alpha$. As we shall see later, this approximation greatly simplifies the differential equations of the two-dimensional stationary model.

Historically, the question of the variation of D_T with V_t is of some interest. In one of the early papers of 1858 Longridge and Brooks³ assumed a velocity squared dependence. In 1875, W. Siemens⁹ attacked this assumption stating that D_T actually varied linearly with V_t . There ensued a debate in which many bitter words but few experimental data were displayed.⁹ In view of our present knowledge that the skin friction force, even in the simplest case of flow past a smooth plate, is the result of complicated boundary layer phenomenon, the existence of this confusion is not surprising.

3.4 Sinking Velocities and Their Relationship to Drag Forces

The studies of submarine cable forces in 1857 and 1858 preceded modern fluid mechanics by many years. To characterize the hydrodynamic forces acting on cable the early investigators used *sinking* or *settling* velocities rather than the more recently conceived drag coefficients. The transverse sinking velocity u_s was defined as the terminal velocity attained by a straight, horizontal length of cable sinking in water.

Similarly, the longitudinal sinking velocity v_s was the terminal velocity of a cable length sinking with its axis constrained to be vertical. If for a given cable the drag forces are functions only of velocity, the parameters w , u_s , and v_s , together with the laws of variation of the drag forces with velocity, completely define the hydrodynamic behavior of the cable. Since sinking velocities are still used in submarine cable technology, it is of interest to relate them to the more modern drag coefficient viewpoint.

In the case of transverse or normal flow around the cable, the variation of D_N with the square of the relative transverse velocity gives $(V_N/u_s)^2 = D_N/w$, since at a transverse velocity equal to the sinking velocity the unit transverse drag force is w . Substituting for D_N from (4) we find

$$u_s = \left(\frac{2w}{C_{D\rho} d} \right)^{\frac{1}{2}} = H. \quad (16)$$

Thus, the transverse sinking velocity u_s is identical with the *hydrodynamic constant* H . We can therefore alternatively write the approximate relationship (9) as

$$\alpha_0 V = u_s, \quad (17)$$

where α_0 is in radians and u_s and V are in knots.

For the tangential or skin friction flow along smooth cable, the sinking velocity concept is inadequate because the unit tangential drag force D_T varies with length as well as with the relative tangential velocity V_t . However, for cable with the conventional jute exterior (cable No. 2), we have $(V_t/v_s)^{1.48} = D_T/w$ and from (15) the vertical sinking velocity v_s is $v_s = (46.1w)^{1/1.48}$, where v_s is in knots.

We note in passing that the cable does not, as is sometimes supposed, sink vertically to the bottom at the transverse sinking velocity u_s . Actually, the term "vertical cable sinking rate" is ambiguous. There are in fact two vertical sinking rates which may be important. Although both these rates are normally approximately equal to u_s neither is identical to it.

Relative to the earth, the resultant velocity V_R of a cable element has two components: a horizontal component of the magnitude of the ship velocity, and a component inclined at the angle α of the magnitude of the cable pay-out rate V_c . These are shown in Fig. 7. The component V_{vert} of V_R , given by $V_{\text{vert}} = V_c \sin \alpha$, is the rate at which a cable *element* sinks vertically. For a laying depth h , the time τ it takes for a cable element to sink to bottom is therefore $\tau = h/V_c \sin \alpha$. This time would, for example, tell one how long it takes a lightweight repeater, integral with the cable, to reach the ocean bottom.

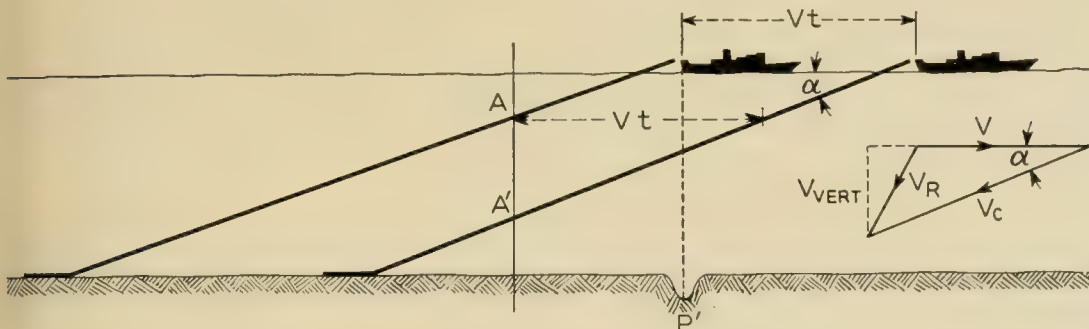


Fig. 7 — Illustration of vertical cable sinking rates.

On the other hand, consider the intersection of the cable *configuration* with a vertical line (Fig. 7). In the time t , as the ship sails a distance Vt , the intersection moves from A to A' , a distance $Vt \tan \alpha$. Hence the cable configuration in this sense sinks vertically at the rate $V \tan \alpha$, and the time θ for the configuration to reach bottom in a depth h is $\theta = h/V \tan \alpha$. One may be interested in how long it takes after the ship has passed over an ocean bottom anomaly P' (Fig. 7) for the cable configuration to reach the anomaly. This is just the time θ .

Hence, the vertical sinking rates $V_c \sin \alpha$ and $V \tan \alpha$ can both be of interest. At the usual ship speeds, $\sin \alpha \approx \tan \alpha \approx \alpha \approx u_s/V$. Further V_c normally differs little from V . Hence, both these rates are indeed normally approximately equal to u_s .

3.5 General Solution of the Stationary Two-Dimensional Model

Assume that each cable element is traveling along the stationary cable configuration with the constant speed V_c . Starting at the ocean bottom let s be the arc length along the stationary configuration. We define s to be positive in the direction opposite to the direction of travel of the cable elements. So, as Fig. 8 indicates, in laying, positive s is directed from the ocean bottom toward the ship, while in recovery the situation is reversed. We let θ be the angle between the positive s direction and the direction of the ship velocity.

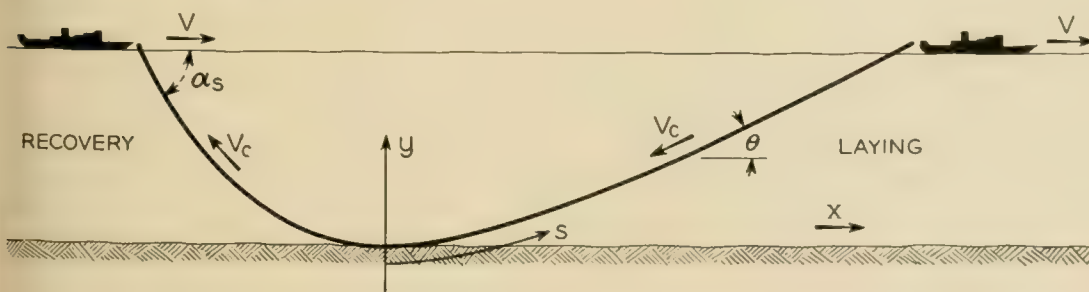


Fig. 8 — Definition of coordinates for the two-dimensional stationary model.

Fig. 9 shows the forces acting on an element of the cable, with tension at the point s being denoted by T . The normal drag force per unit length D_N may, by virtue of (4) and (5), be written in the form

$$D_N = \frac{C_D \rho V^2 d}{2} \sin \theta |\sin \theta|.$$

It is necessary to introduce here $|\sin \theta|$ in order for D_N to have the proper sign for all θ . We note, however, that if $V_c \geq V$, we have from (13)

$$V_t = V_c - V \cos \theta \geq 0.$$

Hence in normal laying and recovery the unit tangential drag force D_T is always in the positive s direction.

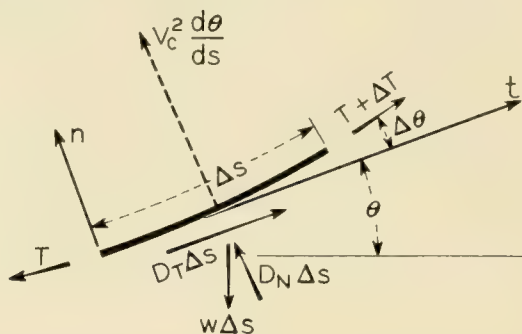


Fig. 9 — Diagram of forces acting on a cable element.

The forces acting on an element produce a centrifugal acceleration $V_c^2 d\theta/ds$. Hence, summing forces along the directions t (tangential) and n (normal), dividing by Δs and sending Δs to zero, we obtain

$$(T - \rho_c V_c^2) \frac{d\theta}{ds} + \frac{C_D \rho V^2 d}{2} \sin \theta |\sin \theta| - w \cos \theta = 0, \quad (a)$$

(18)

$$\frac{dT}{ds} + D_T - w \sin \theta = 0, \quad (b)$$

where ρ_c is the mass density per unit length of cable.

It is seen at the outset that $\theta = \alpha$ is a solution of (18a). It is in fact the important straight-line solution which has been discussed in Section 3.1.

If $\theta \neq \alpha$ and D_T varies only with V_t we may divide (18b) into (18a) and integrate to obtain the solution for T in the following form:

$$\ln \frac{(T - \rho_c V_c^2)}{(T_0 - \rho_c V_c^2)} = \int_{\theta_0}^{\theta} \frac{(w \sin \xi - D_T)}{w(\cos \xi - \Lambda \sin \xi |\sin \xi|)} d\xi, \quad (a)$$

(19)

$$\Lambda = \frac{C_D \rho d V^2}{2w} = \frac{\cos \alpha}{\sin^2 \alpha}, \quad (b)$$

where T_0 is the tension corresponding to the angle θ_0 .

At the cable touchdown point on the ocean bottom only two conditions are possible. If the angle θ is not zero or π there, the cable tension T must be zero. Otherwise a finite tension would act on an infinitesimal length of cable, producing an infinite acceleration. Hence, either the tension T must be zero or the angle θ must be zero or π . The first case normally implies a straight-line configuration (see Appendix A), which has already been discussed. In other cases, we define T_0 as the tension at the touchdown point, and we let θ_0 be zero or π , whichever is appropriate.

If x, y are coordinates in the translating (x, y) frame of a point along the cable configuration, then

$$dx = ds \cos \theta,$$

$$dy = ds \sin \theta.$$

Combining these relations with (18a), we have

$$s = \int_{\theta_0}^{\theta} \frac{(T - \rho_c V_c^2)}{w(\cos \xi - \Lambda \sin \xi | \sin \xi |)} d\xi, \quad (a)$$

$$x = \int_{\theta_0}^{\theta} \frac{(T - \rho_c V_c^2) \cos \xi}{w(\cos \xi - \Lambda \sin \xi | \sin \xi |)} d\xi, \quad (b) \quad (20)$$

$$y = \int_{\theta_0}^{\theta} \frac{(T - \rho_c V_c^2) \sin \xi}{w(\cos \xi - \Lambda \sin \xi | \sin \xi |)} d\xi. \quad (c)$$

Equations (19) and (20) are an integral representation of the complete solution of the basic two-dimensional model. In general, the integrals appearing in these equations cannot be evaluated in terms of elementary functions, and the solution must be obtained by numerical integration. For towing problems where the pay-out velocity is zero, Pode¹² has tabulated these numerical integrations using the approximation that D_T has certain constant values. However, in towing problems the direction of D_T is opposite to what it is in normal laying and recovery problems. Because small magnitudes of D_T were used, these tables nevertheless usually give adequate results in laying and recovery situations as well. At the same time, for submarine cable problems, other approximations allow more convenient ways of evaluating the integrals of (19) and (20).

For example, it is more accurate simply to assume that D_T is zero. As we indicated in Section 3.2, this approximation gives a negligible deviation from the exact solution if the relative tangential velocity V_t is small. Furthermore, in this situation we obtain from (18b)

$$\frac{dT}{ds} = w \sin \theta = w \frac{dy}{ds},$$

and hence the tension at the ship T_s is very nearly

$$T_s = T_0 + wh, \quad (21)$$

where h is the depth at the touchdown point. Thus, if the tangential drag force is negligible, the tension at the ship is essentially the bottom tension plus wh , *regardless* of the nature of the normal drag forces. This is in fact a form of a well-known theorem which, as we shall see in Section 7.1, applies in the three-dimensional case as well.

In the next sections we make further simplifications of the general solution for the specific cases of laying and recovery.

3.6 Approximate Solution for Cable Laying

On long cable lays ship speeds are normally of the order of 4–8 knots, with accompanying values of the critical angle α of the order of 10° – 30° . For these small values of α , the assumption of zero tangential drag together with some mathematical approximations allow further simplifications of the general solution. These simplifications are derived in detail in Appendix C; here we indicate the results. The angle θ which the configuration makes with horizontal is closely given by

$$\tan \frac{\theta}{2} = \tan \frac{\alpha}{2} \left[\frac{1 - [\bar{T}_0/(\bar{T}_0 + \bar{y})]^\gamma}{1 + [\bar{T}_0/(\bar{T}_0 + \bar{y})]^\gamma \tan^4 \frac{\alpha}{2}} \right]^{1/2}, \quad (22)$$

where $\bar{y}$ and $\bar{T}_0$ are dimensionless depth and bottom tension defined by

$$\bar{y} = y/h,$$

$$\bar{T}_0 = T_0/wh.$$

Here we use the cable angle α in the sense of Section 3.2, namely, as a parameter characterizing the hydrodynamic cable properties and the ship speed. The constant γ is in turn defined by

$$\gamma = \frac{(2 - \sin^2 \alpha)}{\sin^2 \alpha}. \quad (23)$$

For small α , $\tan^4 (\alpha/2)$ is negligible and γ is large. Further

$$0 < \frac{\bar{T}_0}{\bar{T}_0 + \bar{y}} < 1.$$

Hence, the denominator in (22) is very nearly unity and θ approaches the critical angle α at small values of $\bar{y}$, even for relatively large values of $\bar{T}_0$ of the order of three or four. Thus in the laying case, the cable configuration is very close to a straight line except for a short distance at the ocean bottom.

In Appendix C it is further shown that for small α

$$\begin{aligned} S &= L + \kappa T_0/w, \\ X &= L \cos \alpha + \lambda T_0/w. \end{aligned} \quad (24)$$

Here S and X are the distance along the cable and the horizontal distances respectively from the touchdown point to the ship (Fig. 11), L and $L \cos \alpha$ are the corresponding distances for straight-line laying at the same ship speed, and κ and λ are functions of the critical angle α which are plotted in Fig. 10. To illustrate the use of (24) we consider the following.

Example: Cable No. 2 is being laid without slack onto a rough bottom from a ship moving at six knots. If the pay-out rate is decreased so the slack is 1 per cent negative, what is the subsequent rise of the tension with time at the ship?

This is really a transient problem. However, we shall try to get an idea of the average behavior of the cable by assuming it passes through a sequence of stationary configurations. Also, we assume that because of the rough bottom there is no slippage of the cable along the ocean floor.

If δ is the amount of negative slack and V the ship speed, then in a time t an amount $V(1 - \delta)t$ of cable will have been paid out. This amount plus the inclined length L will equal the amount contained in the curve AOC (Fig. 11). We then have

$$L + V(1 - \delta)t = S + Vt - (X - L \cos \alpha). \quad (25)$$

Substituting (24) into this equation and solving for T_0 we find $T_0 = (w/(\lambda - \kappa))\delta Vt$ and that by (21) the tension at the ship is given by

$$T_s = wh + \frac{w}{\lambda - \kappa} \delta Vt.$$

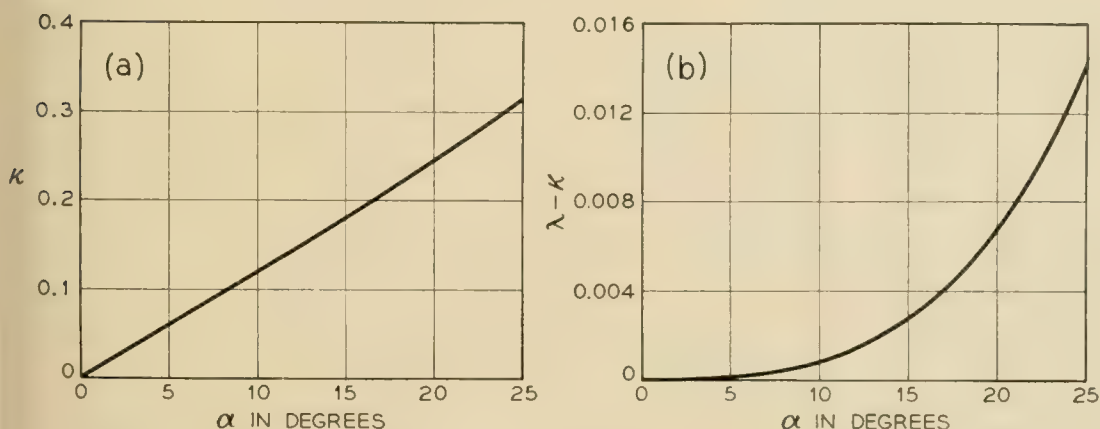


Fig. 10 — Variation of κ and λ with the critical angle.

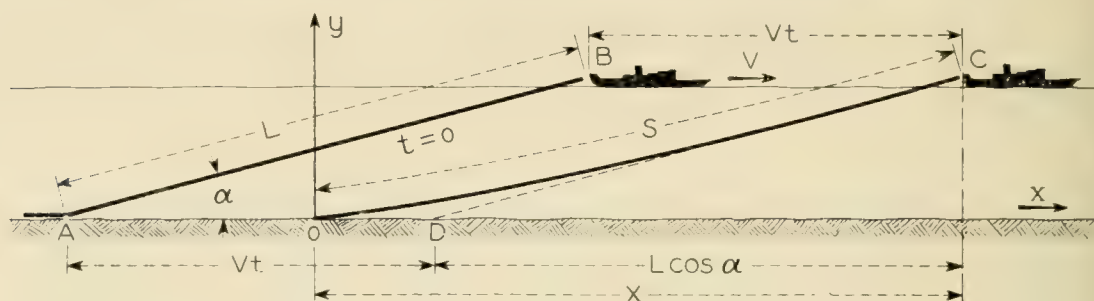


Fig. 11 — Cable geometry at a time t after the onset of negative slack.

For cable No. 2 a ship speed of six knots corresponds to $\alpha = 11.7$ degrees. By Fig. 10, this corresponds to $\lambda - \kappa = 1.4 \times 10^{-3}$. Also $w = 0.705$ lbs/ft by Table I. We get therefore

$$T_s = wh + 3000 \left(\frac{\text{lbs}}{\text{min.}} \right) t.$$

Thus, according to this calculation the tension in this example rises at the extremely rapid rate of 3000 lbs/min. We note also that the rate of tension rise is here independent of the depth h .

In the model which has been postulated, the cable is inextensible; that is, it is assumed not to stretch under load. Because the difference between the lengths of AOC and the sum of the linear segments AD and DC (Fig. 11) is small, one might suspect that the effect of cable extensibility in the present example is important. We can account for this effect in a crude way by assuming that the curve AOC has an additional length corresponding to the stretching caused by the load T_0 acting over the length L . For a cable made of a single material, the stretching would be $T_0 L / EA$, where E is the Young's modulus of the material and A is the cross-sectional area of the cable. In analogy to this we denote the extensile rigidity of the cable by $\overline{EA}$, using the bar to indicate that $\overline{EA}$ is actually a single number obtained directly by measuring the extension of a length of cable loaded in tension. With this notation (25) becomes

$$L + V(1 - \delta)t + \frac{T_0 L}{\overline{EA}} = S + Vt - (X - L \cos \alpha),$$

and repeating the previous computation we find

$$T_s = wh + \left[\frac{wh}{\overline{EA} \sin \alpha} + \lambda - \kappa \right] V \delta t.$$

It is to be noted that in this computation, unlike the inextensible case, the rate of tension rise depends on the depth h .

For conventional helically armored cable, one cannot define a single extensile rigidity because of coupling between pulling and twisting. Thus, how such a cable extends under tension depends on how it is restrained from twisting at the ship and at the ocean bottom. Instead of trying to determine these end restraints, we consider the limiting cases of no restraint and complete restraint to twisting. Data supplied by P. Yeisley indicate the values of EA for cable No. 2 in these conditions to be those given in Table I (Section 3.2). If we take $h = 6,000$ and $12,000$ feet, we find with these values that

$$h = 6000 \text{ feet:}$$

$$\begin{aligned} T_s &= wh + 220 \text{ (lb/min)}t \text{ (twist unrestrained),} \\ &= wh + 640 \text{ (lb/min)}t \text{ (twist restrained),} \end{aligned}$$

$$h = 12,000 \text{ feet:}$$

$$\begin{aligned} T_s &= wh + 120 \text{ (lb/min)}t \text{ (twist unrestrained),} \\ &= wh + 360 \text{ (lb/min)}t \text{ (twist restrained).} \end{aligned}$$

Comparing with the inextensible computation, we see that the extensibility markedly reduces the rate of tension build-up. Nevertheless, even for the case of no restraint to twisting at a depth of $12,000$ feet the rise rate is a relatively high 120 lb/min . Hence, at least over a rough bottom, the tension would quickly indicate the onset of negative slack, although the sensitivity of this indication would decrease with increasing depth.

3.7 Approximate Solution for Cable Recovery

Fig. 8 illustrates how cable is in present practice recovered from the ocean bottom. The cable is in front of the ship as it is brought in over the bow, and the ship pulls itself along the cable. In this process the cable tends to guide or lead the ship directly over its resting place on the ocean bottom.

It is clear that during recovery by this procedure the tension at the ocean bottom is not zero and the cable configuration is not a straight line. Furthermore, in this situation the normal component of the water drag force D_N pushes down on the cable instead of buoying it up as in the case of laying. This in turn implies a higher tension at the ship during recovery than during laying.

If the tangential drag is neglected, the tension at the ship T_s during recovery is given in dimensionless form by (see Appendix C)

$$\frac{\bar{T}_s - 1}{\bar{T}_s} = \left[\tan^2 \alpha \frac{\cos \alpha + \cos \alpha_s}{1 - \cos \alpha \cos \alpha_s} \right]^{1/\gamma}, \quad (26)$$

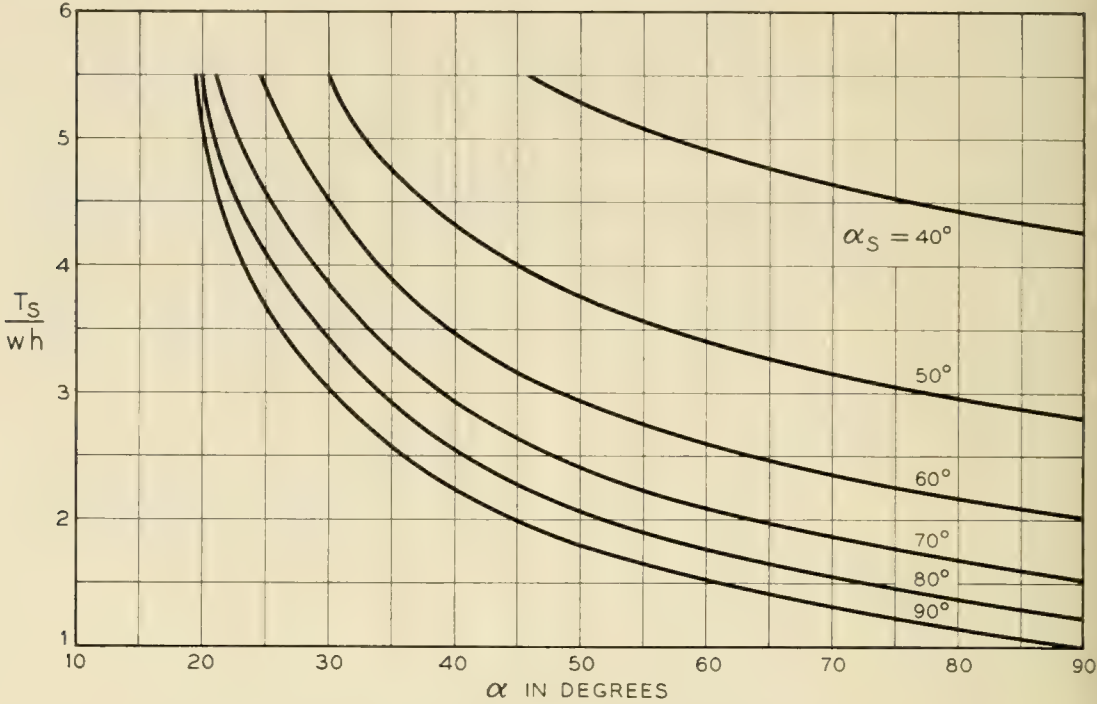


Fig. 12 — Variation of the tension factor for recovery with the critical angle α .

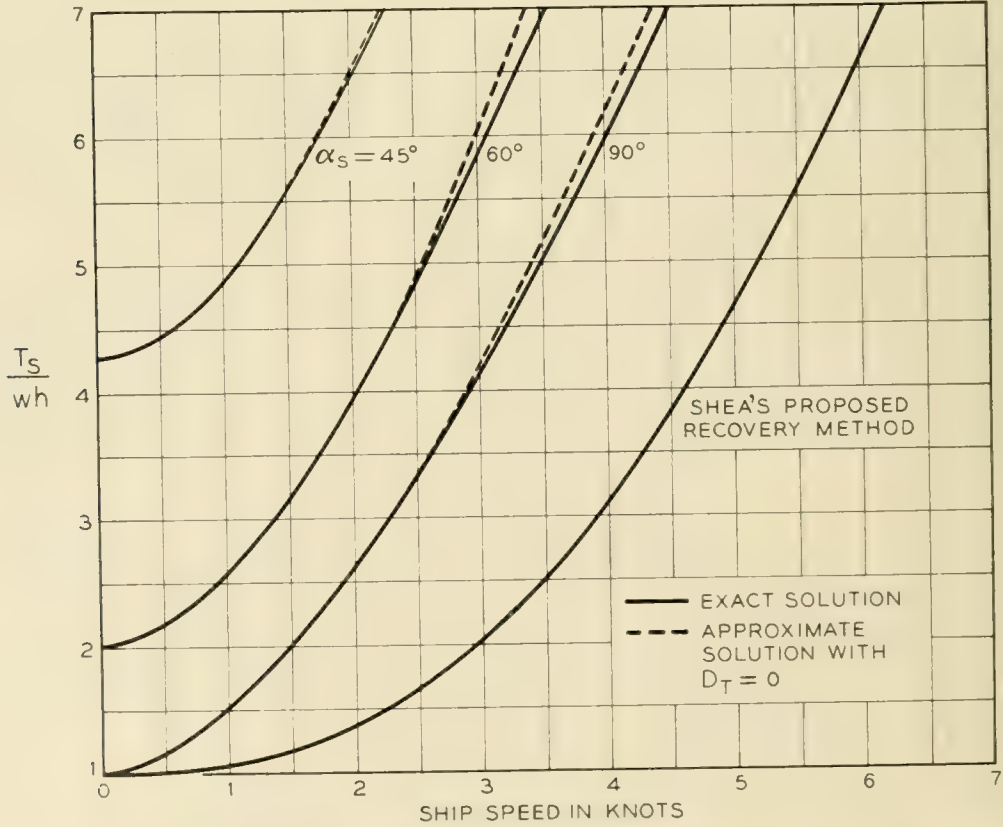


Fig. 13 — Variation of exact and approximate tension factors for recovery of cable No. 2.

where $\bar{T}_s$ is the tension factor defined by $\bar{T}_s = T_s/wh$ and γ is given by (23). Equation (26) is plotted in Fig. 12 in the form of $\bar{T}_s$ versus α for various surface incidence angles α_s (Fig. 8). It is seen that the recovery tensions are in fact considerably higher than the laying tension of approximately wh .

To illustrate the smallness of the error of neglecting the tangential drag force in this computation, we have plotted the approximate and exact curves of $\bar{T}_s$ versus ship velocity for cable No. 2 in Fig. 13. The dotted curves have been computed from (26), while the solid curves have been obtained by substituting D_T from (15) into (19) of the general solution and integrating numerically.* (The curve labeled Shea's recovery method is discussed in the next section.)

The distance along the cable S and the horizontal distance X from the touchdown point to the ship cannot be expressed in a simple form as in the case of laying. However, they can be obtained by numerical integration from (20). The results of this computation for $D_T = 0$ are shown in Figs. 14 and 15.

How Figs. 12, 14 and 15 can be used is illustrated in the following example.

* The standard form of Simpson's rule was used for all the numerical integrations mentioned in the paper. In each case the interval of integration was chosen fine enough to obtain at least three significant figure accuracy.

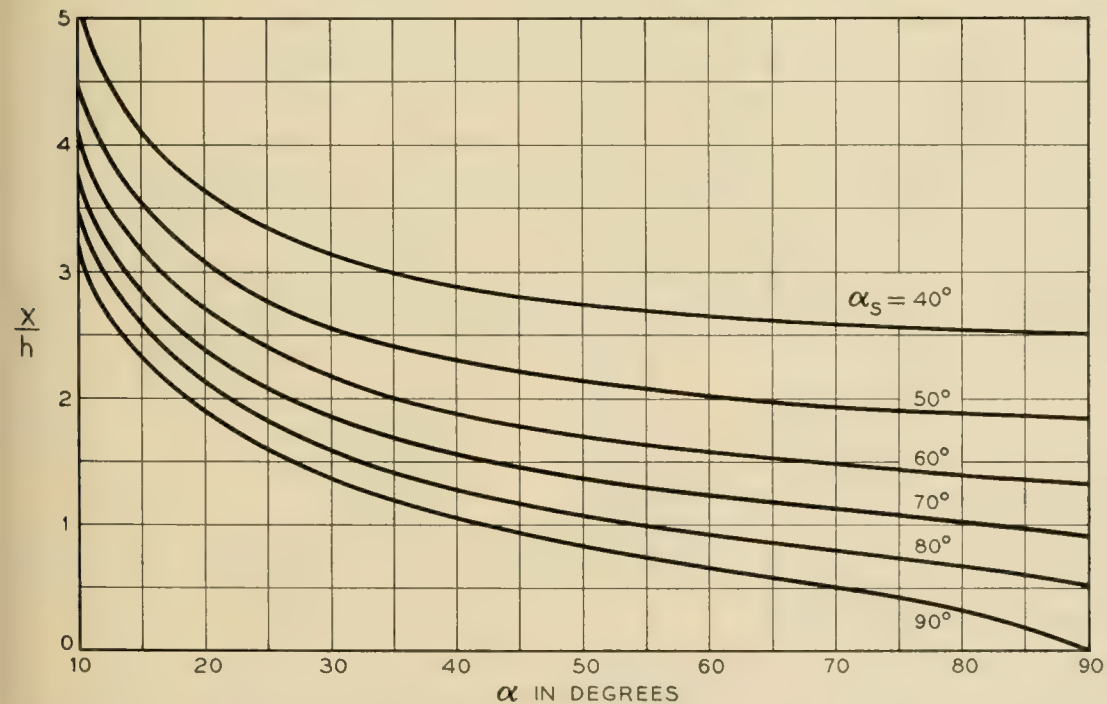


Fig. 14 — Variation of the horizontal distance to the touchdown point during recovery with the critical angle.

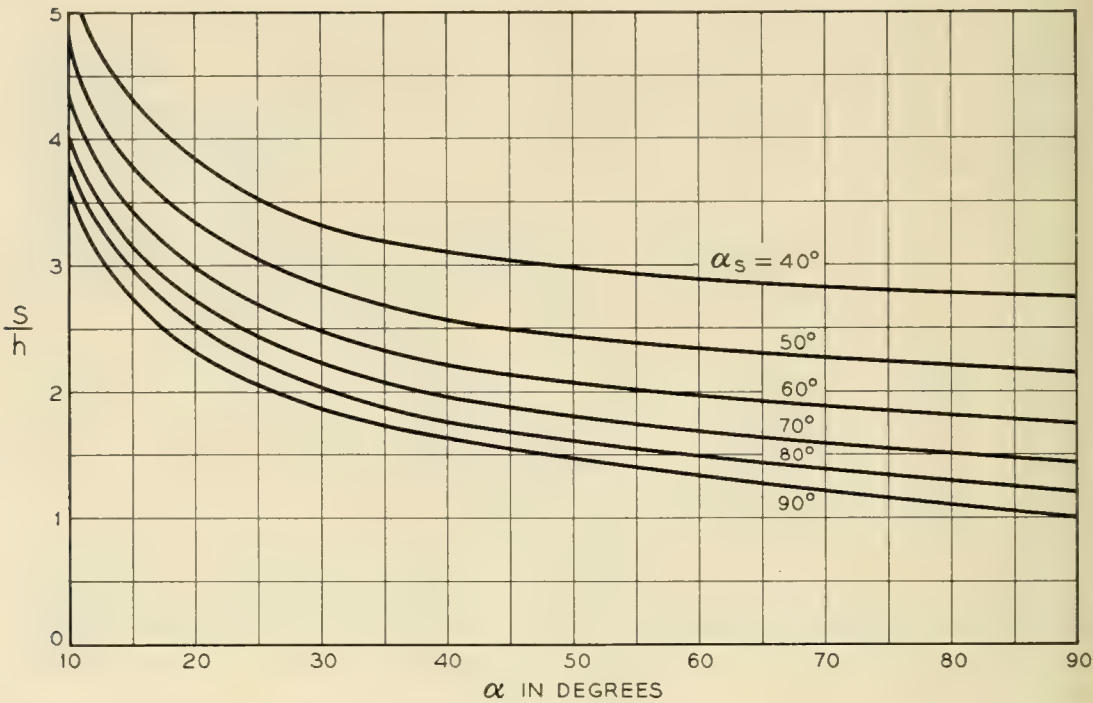


Fig. 15 — Variation of the distance along the cable to the touchdown point during recovery with the critical angle α .

Example: A cable weighing 0.7 lb/ft in sea water and having a hydrodynamic constant H of 70 degree-knots, is to be picked up from a depth of two thousand fathoms. If the ship speed is one knot what is the cable tension at the ship for surface angles α_s of 40° , 60° , and 90° ? How far in front of the ship and how far along the cable will the touchdown point be for these values of α_s ?

As indicated by (9), an H value of 70 degree-knots together with a ship velocity of one knot yields $\alpha_0 = 70$ degrees. Fig. 4 yields in turn $\alpha = 60$ degrees. Entering Fig. 12 with this value of α , we can obtain T_s/wh . In this example the wh tension for a depth of two thousand fathoms is 8,400 lb, and hence the values of T_s/wh and $\bar{T}_s$ are as follows:

α_s	T_s/wh	T_s
40°	4.85	40,700 lbs
60°	2.58	21,700 lbs
90°	1.53	12,900 lbs

From Fig. 14 we get in turn for the horizontal distance from the ship to the touchdown point

α_s	X/h	X
40°	2.65	5300 fathoms
60°	1.56	3120 fathoms
90°	0.66	1320 fathoms

Finally, from Fig. 15 we get for the distance along the cable to the touch-down point

α_s	S/h	S
40°	2.88	5760 fathoms
60°	1.95	3900 fathoms
90°	1.33	2660 fathoms

3.8 *Shea's Alternative Recovery Procedure*

The high tensions which result in the usual recovery operation require slow ship speeds of the order of one knot or less if the cable is not to be broken. One wonders if it is possible to mitigate these tensions and thus speed the recovery process. J. F. Shea discovered that this can theoretically be done by allowing α_s to exceed 90°, thus establishing the straightline configuration (Fig. 16). As in laying, the normal drag forces in this scheme support the cable, rather than push down on it as in conventional recovery. However, in contrast to the laying situation we have $V_t = V_c + V \cos \alpha$. Thus V_t is the sum of V_c and $V \cos \alpha$ instead of their difference and D_T is not necessarily negligible. Furthermore, the direction of D_T is now such as to increase rather than decrease the tension over the wh value. Hence, instead of (2), a summation of forces along the cable yields $T_s = wh + D_T L$, and the tension at the ship can be considerably higher than wh . A curve of T_s as a function of ship speed for cable No. 2 is shown in Fig. 13 with the label "Shea's recovery method". This has been computed for the case of haul-in speed equal to ship speed by means of the above equation and (15). It is seen that the tensions computed for this method of recovery, at least for the cable No. 2, are nevertheless considerably smaller than those which occur in the usual recovery procedure. It would seem that the straight-line recovery technique could fruitfully bear further examination, especially for application to the recovery of long stretches of cable.

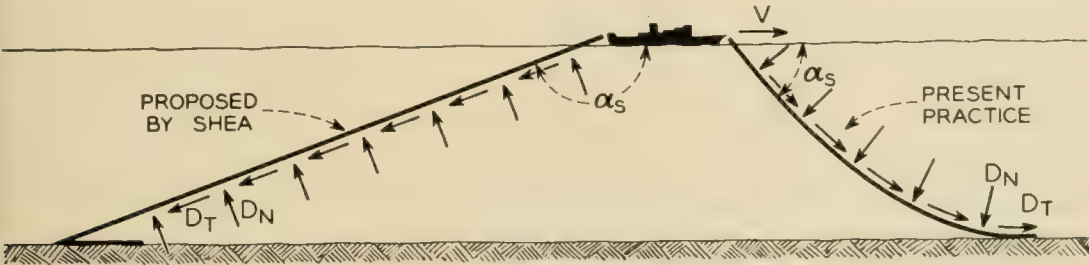


Fig. 16 — The present and Shea recovery methods.

IV. EFFECTS OF SHIP MOTIONS

4.1 *Tensions Caused by Ship Motions*

In the basic stationary model a perfectly calm sea is postulated. However, in reality, wave action gives rise to a random motion of the ship which in turn induces variations in cable tensions around those corresponding to the basic model.

To analyze this effect, we assume that the mean forward velocity of the ship and the mean pay-out or haul-in rate are constant and that the mean tension at the ship and the mean direction of the cable as it enters water are those given by the stationary model. In a reference frame moving with the mean velocity, we resolve the ship displacement into a longitudinal component P_0 (Fig. 17) along the mean or stationary direction and a transverse component Q_0 perpendicular to the stationary direction.

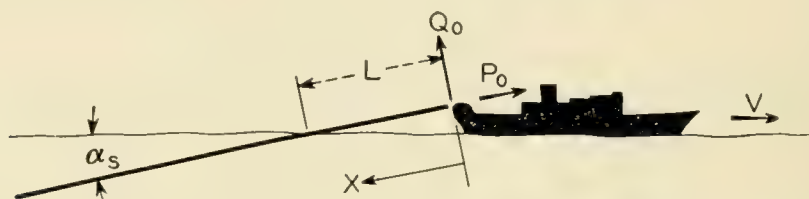


Fig. 17 — Longitudinal and transverse components P_0 and Q_0 of the ship displacement.

Intuitively, one might expect the tensions caused by the transverse displacement Q_0 to be negligible compared to those caused by the longitudinal displacement P_0 . An analysis we have carried through in fact yields this conclusion. Because of its complexity and length, this analysis and the model upon which it is based are given in Appendix D. The results for the case of harmonic variation of Q_0 with time indicate, at least for cable No. 2, that the tension associated with the transverse component Q_0 is indeed negligible for all except ship motions so extreme as to rarely occur.

In addition, this analysis indicates that for the transverse disturbance Q_0 , the amplitude of the responding transverse cable motion decreases exponentially after the cable enters the water because of the damping action of the water drag forces. The “half-life” distance for cable No. 2, that is, the distance along the cable at which the amplitude of a harmonic transverse motion is damped to one-half its surface value, is plotted in Fig. 18 as a function of the period of the motion for various depths h and ship velocities V . The striking feature of these figures is the rapidity of this damping. The analysis thus shows for cable No. 2 that the effect of a transverse disturbance penetrates only a short distance into the water.

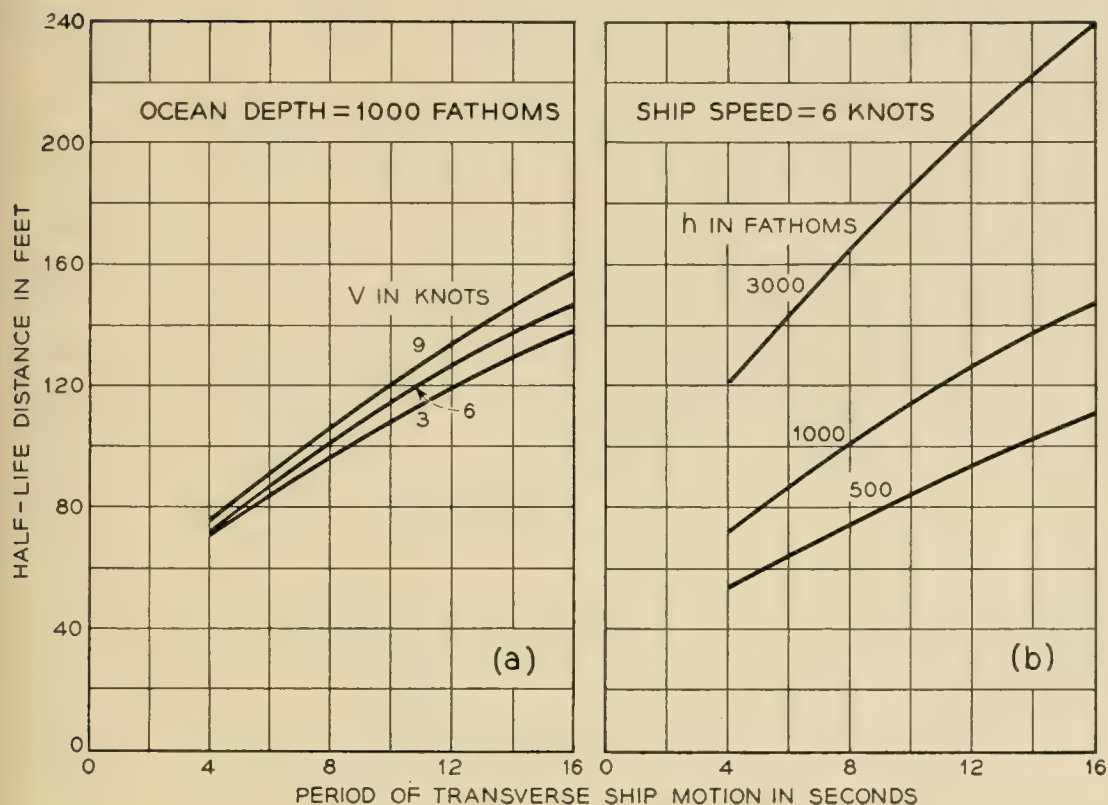


Fig. 18 — Variation of half-life distance of cable No. 2 with the period of ship motion.

As far as cable tensions are concerned, the important ship displacement then is the longitudinal component P_0 , directed along the stationary direction of the cable. The analysis of Appendix *D* leads to the basic one-dimensional wave equation

$$\frac{\partial^2 p}{\partial x^2} - \frac{1}{c_1^2} \frac{\partial^2 p}{\partial t^2} = 0 \quad (27)$$

for the description of the longitudinal motion. In this equation p is the deviation in longitudinal displacement from the mean pay-out or haul-in displacement, and the remaining symbols are defined as (Fig. 17)

x = distance from the mean ship position along the stationary cable configuration,

t = time,

$c_1^2 = \overline{EA}/\rho_c$.

The additional tension T_p due to ship motion is in turn given by

$$T_p = \overline{EA} \frac{\partial p}{\partial x}. \quad (28)$$

As in the example of Section 3.6, we have again assumed that by using

in (28) the limiting values of $\overline{EA}$ obtained by complete restraint to twisting and no restraint to twisting during pulling, one can obtain bounds on the actual displacements and tensions.

The solution of (27) under arbitrary boundary conditions can be obtained from standard textbooks. Probably it is most representative to assume the cable is semi-infinite. That is, although damping of the cable is normally so small that we neglect it in (27), we may assume, because of the cable's great length, that the damping is sufficient to cause complete decay of a disturbance initiated at the ship, and that such a disturbance is not reflected from the ocean bottom. Under this condition the additional tension T_p is given by

$$T_p = -\sqrt{EA\rho_c} \frac{dP}{dt}, \quad (29)$$

where $P = P_0 + P_1$ with dP_1/dt being in turn the increment in pay-out rate or decrement in haul-in rate from the mean. For cable No. 2, Table I (Section 3.2) indicates that

$$\begin{aligned} \sqrt{EA\rho_c} &= 220 \text{ lb/ft/sec (twist unrestrained),} \\ &= 400 \text{ lb/ft/sec (twist restrained).} \end{aligned}$$

Two examples will make clear the application of (29).

Example 1: Steady-State Laying or Recovery in a Regular Seaway.

Assume that in a frame of reference traveling at the mean horizontal ship velocity ship surging (to and fro forward motion) is zero and the combined heave and pitch motion is normal to the ocean surface and is given by

$$W = A \sin 2\pi \frac{t}{\tau}.$$

If the period τ is 6 seconds and the amplitude A is 15 feet find for cable No. 2,

- a) $(T_p)_{\max}$ for laying at a constant pay-out rate and at a ship speed of 6 knots,
- b) $(T_p)_{\max}$ for recovery at a constant haul-in rate and with a surface incidence angle of 60° .

In both cases (a) and (b), the deviation P_1 in pay-out or haul-in rate is zero, hence $P = P_0 = W \sin \alpha_s$, and $(dP/dt)_{\max} = (2\pi/\tau) A \sin \alpha_s$. Since

cable No. 2 has a hydrodynamic constant H of 70 degree-knots, we have in case (a)

$$\alpha_s = 11.7 \text{ degrees and } \left(\frac{dP}{dt} \right)_{\max} = 2.12 \text{ ft/sec.}$$

From (29) we get therefore

$$\begin{aligned} (T_p)_{\max} &= 466 \text{ lbs (twist unrestrained),} \\ &= 848 \text{ lbs (twist restrained).} \end{aligned}$$

In case (b) we have $\alpha_s = 60^\circ$, $(dP/dt)_{\max} = 13.6 \text{ ft/sec}$, and hence

$$\begin{aligned} (T_p)_{\max} &= 2,990 \text{ lbs (twist unrestrained),} \\ &= 5,430 \text{ lbs (twist restrained).} \end{aligned}$$

During recovery by conventional methods the surface incidence angle α_s is in general much larger than that which occurs during laying. The above example points up that one can expect correspondingly larger ship motion tensions during recovery than during laying in the same sort of seas. Since the stationary tensions are also much larger during recovery, recovery is the condition for which the strength of the cable should be designed.

In this example we have considered a regular seaway, something which does not exist in nature. Recent work in the application of the theory of stochastic processes to the study of ocean waves and ship dynamics promises to develop into a realistic description of the behavior of ships at sea.¹³ When such a description becomes available, we shall be able to obtain a better estimate of the magnitudes of ship motion tensions.

As far as data presently available are concerned, the maximum storm condition vertical velocity at the bow or stern recorded by the *U.S.S. San Francisco* during her research voyage of 1934 was 22 feet/sec.¹⁴ Since this ship was roughly the size of a cable ship such as the *H.M.S. Monarch*, this figure might indicate the order of the maximum velocities to be expected in cable practice. In terms of our example, for six knot laying this vertical velocity would imply

$$\begin{aligned} T_p &= 980 \text{ lbs (twist unrestrained),} \\ &= 1,780 \text{ lbs (twist restrained).} \end{aligned}$$

For recovery at a surface incidence angle of 60° , it would imply in turn

$$\begin{aligned} T_p &= 4,200 \text{ lb (twist unrestrained),} \\ &= 7,600 \text{ lb (twist restrained).} \end{aligned}$$

However, it is to be cautioned that these numbers are merely indicative and might differ considerably from those which occur on a particular cable ship.

Example 2: Brake Seizure

While laying cable No. 2 at six knots in a perfectly calm sea, a sudden seizure of the brake occurs. What is the resulting initial rise in tension?

Because of the calm sea we have $P_0 = 0$. Therefore

$$\frac{dP}{dt} = P_1 = -V \cos \alpha_s.$$

With the value of $V = 6$ knots and a corresponding α_s of 11.7° (see Example 1) we have $dP/dt = 9.9$ ft/sec and hence from (29)

$$\begin{aligned} T_p &= 2180 \text{ lbs (twist unrestrained),} \\ &= 3970 \text{ lbs (twist restrained).} \end{aligned}$$

These values of T_p pertain only to the transient values occurring while the tension wave is being transmitted to the ocean bottom. If the seizure in this case occurred at a depth of three thousand fathoms, the time of transit to the ocean bottom would be only of the order of nine seconds. After reaching bottom our initial assumption of no reflection from the bottom would be violated and (29) would no longer hold. In reality the cable tension would continually increase at the ship and reversing ship engines or some other action would be required to avoid rupture of the cable.

V. DEVIATIONS FROM A HORIZONTAL BOTTOM

5.1 *Kinematics of Laying Over a Bottom of Varying Depth*

Ocean bottom topography is not everywhere flat and horizontal as postulated in the basic model. In the Mid-Atlantic ridge, for example, there exist bottom slopes of thirty or forty degrees. In other places submarine canyons with almost vertical sides have been found. Furthermore, where the bottom is steepest it is most likely to be rocky and craggy since erosion tends to smooth out a sandy or muddy bottom. Therefore, it is important to know how the cable should be paid out to cover a bottom of varying depth. To help determine this, we extend here the stationary model to the case of a non-horizontal bottom.

In Section 3.1 we indicated that if the cable tension at the touchdown point is zero the configuration according to the basic model is a straight line, regardless of how the cable is paid out. If the cable is paid out with

slack with respect to the bottom, the zero touchdown tension condition is fulfilled. Hence, under the proper slack pay-out, the cable geometry and, as we shall see, the cable kinematics are particularly simple.

Essentially, we must consider two deviations from the horizontal bottom, namely, downhill or descent laying and uphill or ascent laying. We consider these situations in turn, confining ourselves to bottoms of constant slope since any bottom contour can be approximated by straight-line segments.

To cover a descending bottom, the cable pay-out rate must exceed the ship speed, Fig. 19(a). To cover an ascending bottom, the angle of incidence α of the cable, which as we have seen in Sections 3.1 and 3.2 depends only on the ship speed, must exceed the ascent angle γ , Fig. 19(b). Otherwise, the situation shown in Fig. 19(c) develops. Hence the critical parameters are pay-out speed and ship speed.

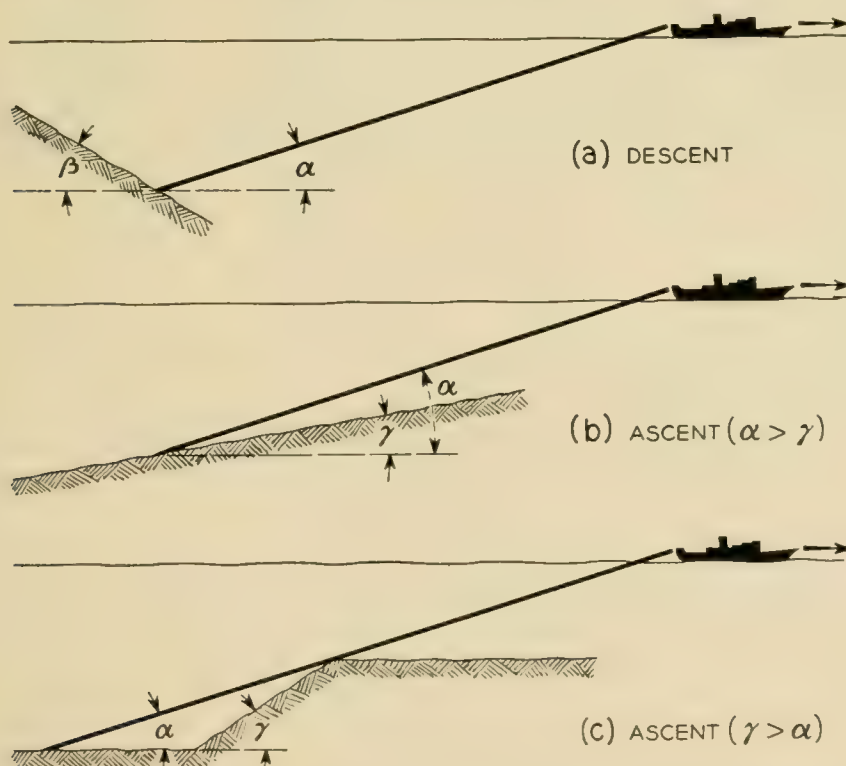


Fig. 19 — Cable geometry during straight-line descent and ascent laying.

During descent laying we see from Fig. 20 that in a time t an amount of cable equal to $a + b$ must be paid out. Hence the required pay-out rate V_c is $(a + b)/t$. But by straightforward trigonometry

$$V_c = \frac{a + b}{t} = \frac{\sin \alpha + \sin \beta}{\sin (\alpha + \beta)} V, \quad (30)$$

where β is the angle of descent and α is the straight-line incidence angle.

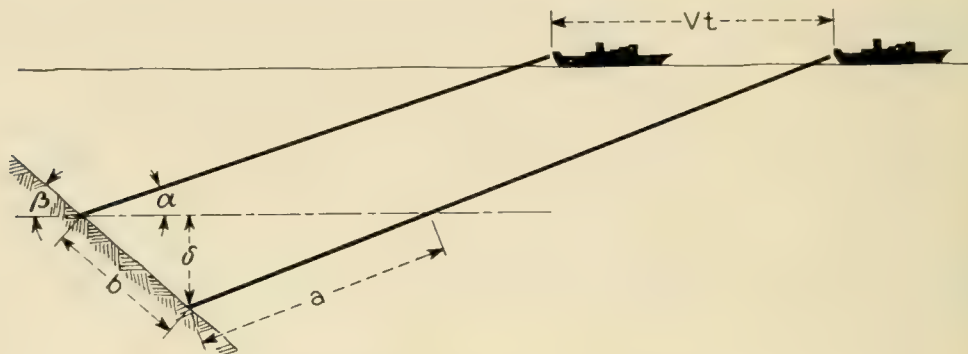


Fig. 20 — Kinematics of straight-line descent laying.

In accordance with usual terminology we define the *slack* ϵ as

$$\epsilon = (V_c - V)/V. \quad (31)$$

We shall think of the slack as being composed of two parts: a *fill* f , which is the amount of slack required for the cable to cover the bottom, and an *excess*, equal to $\epsilon - f$, which will normally be laid to provide a margin of safety. Substituting V_c from (31) into (30) we get the expression for the fill f . The result can be transformed to the form

$$\tan \frac{\alpha}{2} \tan \frac{\beta}{2} = \frac{f}{2 + f}. \quad (32)$$

The quantities α , β and f are normally all small quantities and we may make the approximations

$$\tan \frac{\alpha}{2} \approx \frac{\alpha}{2},$$

$$\tan \frac{\beta}{2} \approx \frac{\beta}{2},$$

$$\frac{f}{2 + f} \approx \frac{f}{2}.$$

For $\alpha, \beta < 30^\circ$ and $f < 0.06$, the error in each of these approximations is less than 3 per cent. Hence, with good accuracy we write (32) as

$$f = \frac{\alpha\beta}{2}, \quad (33)$$

where α and β are expressed in radians. Further, we have from (9) that α in radians is very nearly $\alpha = H/V$, where H is in radian knots. Substituting this expression into (33), we get for the fill

$$f = \frac{H\beta}{2V}. \quad (34)$$

Finally, using this expression for f in (31), we arrive at*

$$V_c - V = \frac{H\beta}{2}. \quad (35)$$

Thus, the *increment* in required pay-out rate is essentially a function only of the descent angle β and is *independent* of ship speed.

In the case of an ascending bottom for which $\alpha > \gamma$, Fig. 19(b), positive bottom slack may be obtained with a pay-out of less than the ship speed. The allowable decrement in pay-out rate is given by

$$V - V_c = \frac{H\gamma}{2}, \quad (36)$$

that is, the same as the required increment for ascent laying. Likewise, the fill f in this case is simply $f = -(H\gamma/2V)$.

The only way to avoid the situation shown in Fig. 19(c) where $\alpha < \gamma$ is to sail slowly enough to maintain an incidence angle α greater than the angle of rise γ . By (9), we have for most laying speeds $\alpha V \approx H$. With good accuracy the condition $\alpha > \gamma$ thus implies

$$V < \frac{H}{\gamma}. \quad (37)$$

Therefore, for a given rise γ the limiting ship speed is simply H/γ .

5.2 Time-Wise Variation of the Mean Tension in Laying Over a Bottom of Varying Depth

In the cases where the cable is paid out with excess onto a bottom of constant slope, the variation of the mean tension at the ship with time is easily computed. During descent laying the increase in depth δ after a time t is by elementary trigonometry (Fig. 20)

$$\delta = \frac{\sin \alpha \sin \beta}{\sin (\alpha + \beta)} Vt.$$

Hence, the rate of rise of the mean shipboard tension is

$$\frac{dT}{dt} = \frac{w\delta}{t} = \frac{\sin \alpha \sin \beta}{\sin (\alpha + \beta)} wV. \quad (38)$$

Similarly, during an ascent lay for which the bottom is less steeply

* Note that in (33), (34) and (35), H may be replaced by the numerically identical transverse settling velocity u_s (see Section 3.4).

inclined than the cable ($\alpha > \gamma$), the rate of decrease in shipboard tension is

$$\frac{dT}{dt} = -\frac{\sin \alpha \sin \gamma}{\sin (\alpha - \gamma)} wV.$$

Like negative slack laying on a flat bottom, the variation of tension with time depends greatly on the frictional characteristics of the bottom in cases other than the above. We therefore limit ourselves to situations where the cable does not move with respect to the ocean floor. This case might be approximated by rough bottoms, where the cable might wedge itself between rocks.

A nomograph giving a rough estimate of the rise of mean tension with time when a cable becomes completely suspended is worked out in Appendix E. In deriving this nomograph it is assumed that the cable takes on a sequence of stationary configurations. This assumption is probably reasonable if the time span of the tension rise is large compared to the time of passage of a tension wave from the ship to ocean floor and return, which as mentioned in Section 4.1 is of the order of 18 seconds. However, because of this assumption and others mentioned in Appendix E, we regard the tension variation computed by the nomograph only as a crude approximation.

Fig. 21 shows the mean ship-board tension versus time computed by means of the nomograph for various slacks ϵ , where ϵ is defined by (31). The values which were used for the other parameters entering the calculation were

$$\frac{wh}{EA} = 3.1 \times 10^{-3},$$

$$\alpha = 12^\circ.$$

Also shown on this curve is the tension rise computed for the case of laying down a vertical slope without excess. The rise for this case is given by (38) with $\beta = 90^\circ$. It is seen that as the slack ϵ is increased the curves for a complete suspension approach the $\beta = 90^\circ$ curve. Indeed, it can be shown that under the assumptions made in computing Fig. 21 the $\beta = 90^\circ$ curve gives a lower bound on the tension rise with time in the case of a complete suspension. A tension rise rate greater than the $\beta = 90^\circ$ rate is thus an indication of unsatisfactory covering of the bottom.

In the case of too rapid a ship speed resulting in $\alpha < \gamma$ (Fig. 19c) restraint of movement of the cable along a rough bottom would cause the tension on the high side of the crest to be zero. There would thus be a sudden drop in tension corresponding to the sudden decrease in depth at the touchdown point after the cable was laid over the crest of the hill.

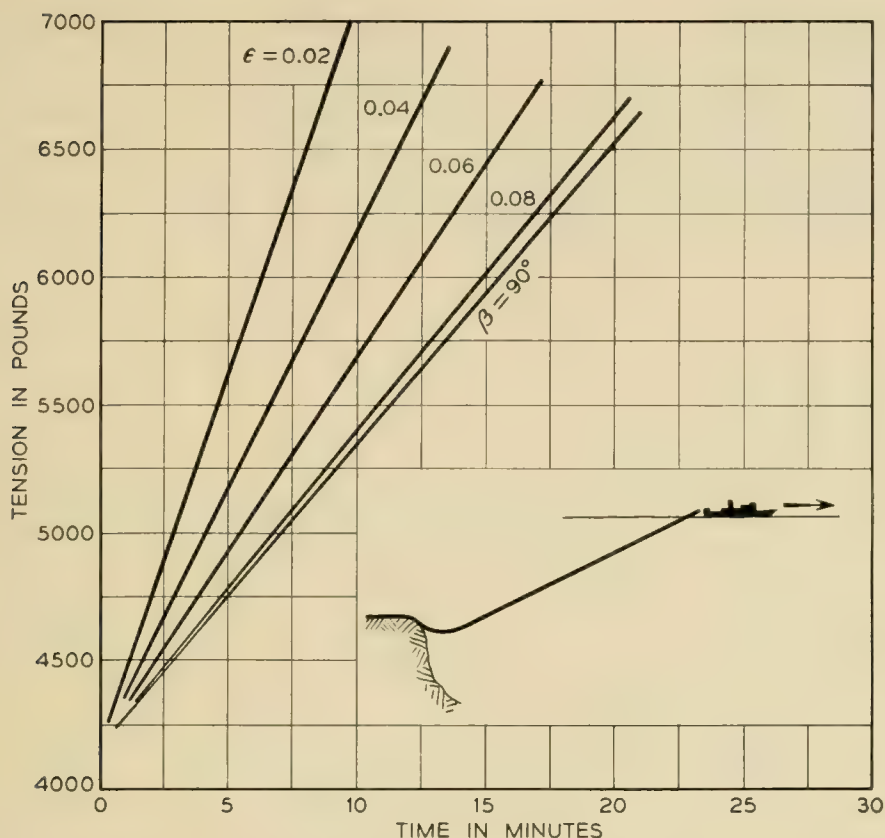


Fig. 21 — Variation of tension with time when cable No. 2 is completely suspended.

On the other hand, in the case of a frictionless bottom, the removal of the supporting water drag forces would cause the cable to seek a catenary equilibrium position on the low side of the crest. But in doing this, the cable would drag itself over the crest, with an accompanying *increase* in shipboard tension.

Thus for the case of a bottom rise steeper than the cable inclination ($\alpha < \gamma$) either an increase or a decrease of tension with time is possible, depending on the nature of the bottom.

5.3 Residual Suspensions

If the cable is not paid out rapidly enough, or if the ship speed is excessive, the cable will be left with residual suspensions after it has been laid. To get an idea of the possible magnitudes of the tensions accompanying these suspensions, we consider here some numerical examples pertaining to cable No. 2. As before, we assume for definiteness the extreme case of a bottom rough enough to prevent movement of the cable.

In Fig. 22 is shown the profile of a 35 fathom (210 feet) increase in depth with a maximum slope of 45° . This profile was obtained from

fathometer records of the Mid-Atlantic Ridge provided by Professor Bruce C. Heezen of the Lamont Geological Observatory. Laying down this slope at a ship speed of six knots requires a slack of 8.5 per cent (see Section 5.1). If the slack were only 5 per cent, the successive cable configurations as calculated by the methods of Appendix E would be those shown in Fig. 22(a).^{*} The cable would touch bottom after 2.6 minutes,

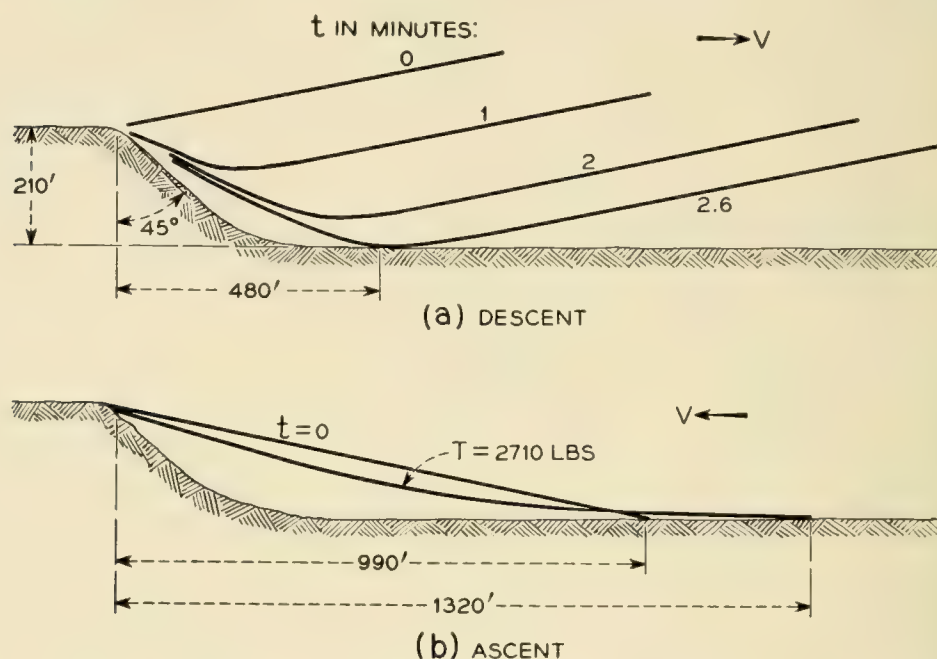


Fig. 22 — Successive cable configurations during a 35-fathom descent and ascent lay of cable No. 2 at a 6,000-ft depth with an assumed 5 per cent slack and 6-knot ship speed.

leaving a residual suspension with a half-span of 480 feet and a tension of 525 lbs. The mean tension at the ship would correspondingly increase by 525 lbs during the 2.6 minute time interval. In even moderately rough seas, this tension change could be obscured by the ship motion tensions.

Consider this profile next to represent an ascending lay under a ship speed of six knots. Fig. 22(b) shows the initial ($t = 0$) and residual cable configuration. Because of the small incidence angle of the initial straight-line shape, the residual half-span of the catenary is a quarter of a mile (1320 feet) long, and the accompanying residual tension is 2,710 lbs, or roughly that which normally occurs in laying at a depth of $\frac{3}{4}$ of a nautical mile. At the ship, there would be a decrease in the mean tension of 130 lbs. corresponding to the 35 fathom decrease in depth. Again, a tension change of this magnitude would be difficult to discern because of ship motion tensions.

^{*} We have further taken the ratio $wh / \overline{EA}$ to be 3.1×10^{-3} in this computation. However, the results are very insensitive to change in the $wh / \overline{EA}$ ratio.

If the above 35 fathom change occurred at a depth of say three thousand fathoms, a very sensitive fathometer would be required to detect it. Thus, although complete restraint of cable movement along the bottom is an extreme and unlikely condition, the above examples indicate that long residual suspensions can occur with essentially no manifestation at the ship, especially in deep water.

VI. CABLE LAYING CONTROL

6.1 General

We have seen that the mean cable tension at the ship reflects the amount of slack which is being paid out and how the cable is covering the bottom. However, in most cases this reflection is not sensitive. For example, the tangential drag force D_T varies with V_t , the longitudinal velocity of the cable relative to the water. In theory, as (2) shows, one can therefore determine the amount of slack being paid out from ship-board tension measurements. For cable No. 2, we have plotted in Fig. 23 the variation of the mean tension at the ship as a function of slack for a ship speed of six knots and a depth of two thousand fathoms. At three per cent slack the tension is 8,240 pounds, while at six per cent slack it is 8,020 pounds, a difference of only 220 pounds. This amount of tension

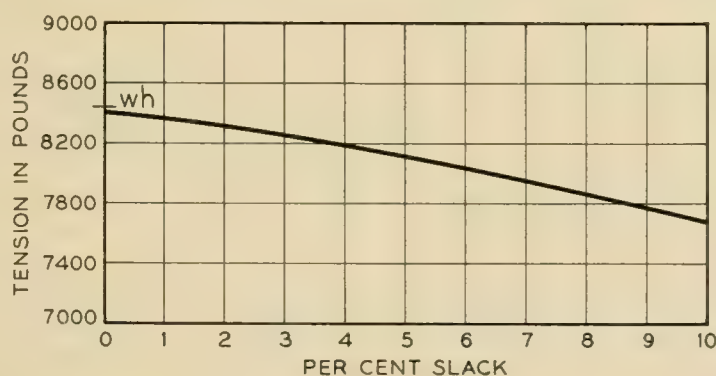


Fig. 23 — Variation of shipboard tension with per cent slack for laying cable No. 2 at a ship speed of 6-knots in a depth of two thousand fathoms.

could be easily obscured by the effect of ship motion. Thus, to measure slack accurately by relating it to cable tensions one would have to know the depth and cable parameters very precisely and, in addition, would need a very efficient filter to separate out the “noise” tension caused by ship motion.

Similarly, it has been shown that residual suspensions can occur with essentially no reflection in the tension readings at the ship. Hence, although tension readings can give a valuable check on how the cable is

covering the bottom, it would seem difficult for them to provide exact enough data for the control of cable laying.

At the same time we have seen that if the bottom contour is known in advance, then for a given ship speed one can compute the required cable pay-out rate. Also with foreknowledge of the bottom, one can anticipate steep bottom ascents and decrease the ship speed accordingly. Such a purely kinematic attack on the cable laying problem would seem more fruitful than an attack which depends on measurements of ship-board cable tensions.

Possibly the simplest way of measuring the bottom contour is by means of a fathometer located at the ship. Since the cable ship is normally far forward of the touchdown point of the cable, one could in theory obtain in this manner the required advance knowledge of the contour. In present practice, a taut piano wire is used to obtain the ground speed of the ship. We examine briefly the accuracy of this method in the next section.

6.2 Accuracy of the Piano Wire Technique

The taut wire is laid simultaneously with the cable, but under a constant mean shipboard tension. If the bottom is perfectly horizontal, the speed of the wire coincides with the ground speed of the ship. However, when the bottom depth is variable and the wire is laid up and down hill, the wire's pay-out speed deviates from the ship speed. By (31), it is seen that the error in the ship speed which is indicated by the wire is just equal to the slack ϵ with which the wire is paid out. This slack, which may be positive or negative, can in turn be estimated by the methods of the previous sections.

Consider the beginning (denoted by (1) in Fig. 24) and end (denoted by (2) in Fig. 24) of a downhill lay of the piano wire. As before we neglect the tangential drag force. Then, the condition that the tension at the ship remains constant gives, by (21),

$$(T_0)_1 + wh_1 = (T_0)_2 + wh_2, \quad (39)$$

where the subscripts 1 and 2 refer to the configurations at the beginning and end of the downhill lay. If $\bar{\epsilon}$ is the average slack or error of the piano wire during the descent, then $(1 + \bar{\epsilon})V$ is its average pay-out rate, and we have by Fig. 24,

$$S_1 + (1 + \bar{\epsilon})Vt = \frac{h_2 - h_1}{\sin \beta} + S_2, \quad (40)$$

where S_1 and S_2 are lengths along the cable from the touchdown point to the ship. Also, from Fig. 24

$$X_1 + Vt = (h_2 - h_1)/\tan \beta + X_2. \quad (41)$$

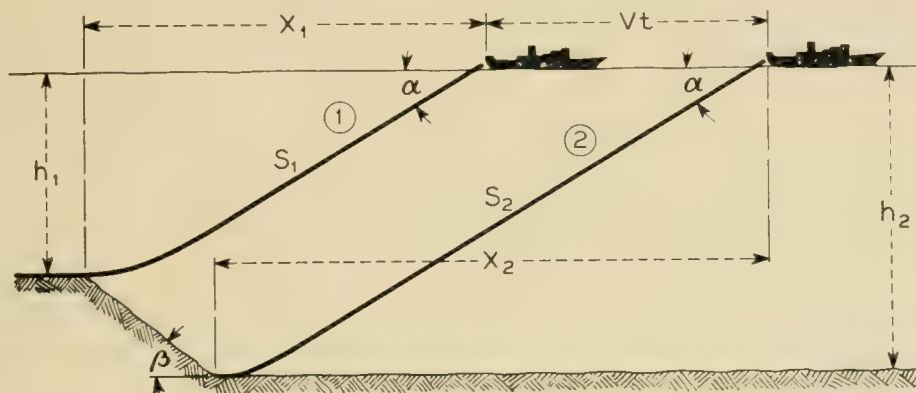


Fig. 24 — Piano wire configurations at the beginning and end of a descent lay.

Equations (39), (40), and (41), together with the general equations (Section 3.6)

$$\begin{aligned} S &= \frac{h}{\sin \alpha} + \kappa \frac{T_0}{w}, \\ X &= \frac{h}{\tan \alpha} + \lambda \frac{T_0}{w}, \end{aligned} \quad (24)$$

allow one readily to solve for the average error $\bar{\epsilon}$ in the piano wire indication of ship ground speed. The result is

$$\bar{\epsilon} = \frac{\sin \alpha + \sin \beta - \kappa \sin \alpha \sin \beta}{\sin (\alpha + \beta) - \lambda \sin \alpha \sin \beta} - 1. \quad (42)$$

For the small values of α and β which normally occur during the laying of the piano wire, the terms $\lambda \sin \alpha \sin \beta$ and $\kappa \sin \alpha \sin \beta$ are negligible. Hence, the average error $\bar{\epsilon}$ is thus very nearly

$$\bar{\epsilon} = \frac{\sin \alpha + \sin \beta}{\sin (\alpha + \beta)} - 1, \quad (43)$$

which, as (30) indicates, coincides with the amount of fill which would be required to lay downhill with the straight-line or zero touchdown tension configuration. Equation (43) is in turn closely approximated by (Section 5.1)

$$\bar{\epsilon} = \frac{\alpha\beta}{2} = \frac{H\beta}{2V}. \quad (44)$$

Similarly, for ascent laying of the wire on a bottom which rises less steeply than the inclination of the wire (19b), we get

$$\bar{\epsilon} = \frac{\sin \alpha - \sin \gamma + \kappa \sin \alpha \sin \gamma}{\sin (\alpha - \gamma) + \lambda \sin \alpha \sin \gamma} - 1, \quad (45)$$

which is very nearly

$$\bar{\epsilon} = \frac{-\alpha\gamma}{2} = -\frac{H\gamma}{2V}. \quad (46)$$

Thus, in both the above cases, the error in the piano wire technique can be closely obtained by assuming that the configuration of the wire is a straight line during ascent and descent laying. This is not surprising since, as we saw in Section 3.6, the deviation from the straight-line configuration during piano wire laying is normally small.

Because of its smooth exterior, the normal or transverse drag coefficient of the piano wire probably can be obtained from published curves for flow past a smooth right circular cylinder as shown in Appendix B. For typical 12 gauge (0.0290 inch diameter) piano wire, these curves yield a value of C_D of 1.45 and an H value of 25.0 degree-knots. However, these values of C_D and H must be considered tentative until confirmed experimentally.

Knowing the wire's H value, we can compute the error of the ground speed caused by descent and ascent laying of the piano wire by means of (44) and (46). The result of this computation for $H = 25.0$ degree-knots is shown in Fig. 25.

When the ascent angle of the bottom exceeds the incidence angle of the wire, suspensions result and the error cannot be computed without

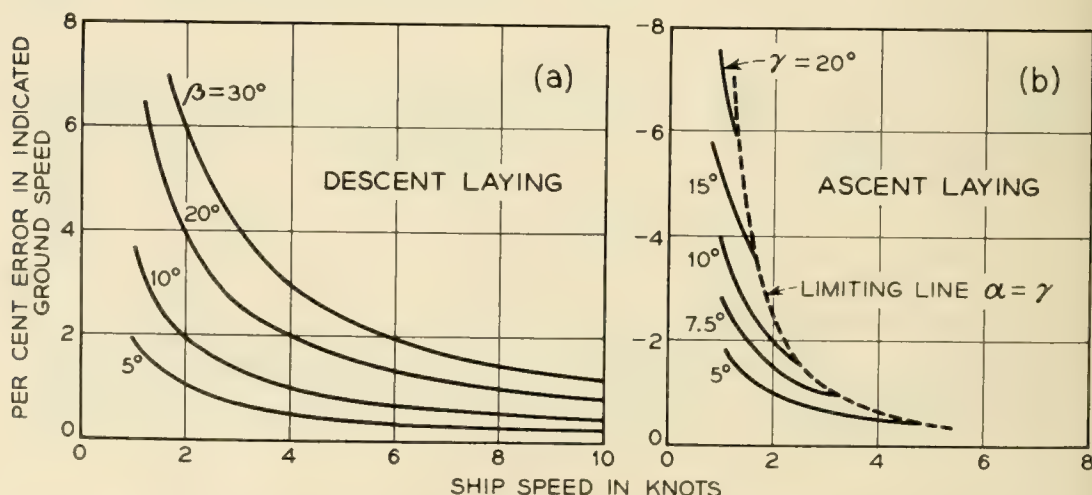


Fig. 25 — Error during descent and ascent laying of 12-gauge piano wire.

knowledge of the frictional properties of the bottom. For an H value of 25.0 degree-knots, (37) indicates that suspensions will occur for ship speeds V greater than $V = 25.0/\gamma$, where V is in knots and the ascent angle γ is in degrees. Hence, for a typical laying speed of 6 knots, ascent angles greater than 4.2 degrees will cause suspensions of the piano wire. These magnitudes indicate that suspensions of the piano wire probably actually develop in practice.

It is seen from Fig. 25 that for the usual small ascent or descent angles, the piano-wire technique is quite accurate, while for large bottom slopes it can be considerably in error. Again, however, if the bottom contour is known in advance, these errors can be estimated in the cases plotted in Fig. 25 and therefore can be corrected for. In this manner, the piano wire could be improved to give accurate ground speeds in all two-dimensional situations, with the exception of the case of a suspension caused by a too steeply ascending bottom. Such suspensions can be avoided only by maintaining a sufficiently slow ship speed. However, as seen by the small computed H value of 25.0 degree-knots, the ship speeds required to avoid piano wire suspensions on uneven bottoms are probably prohibitively slow. Hence, for steeply ascending bottoms it is likely that some other means of determining the ship ground speed is necessary.

VII. THREE-DIMENSIONAL STATIONARY MODEL

7.1 General

Thus far we have assumed that the cable lies entirely in the plane formed by the ship's velocity vector and the gravity vector. Because of the symmetry of the cable cross-section, this assumption seems reasonable.* However in certain cases, as for example in the presence of ocean cross-currents, the assumption of a planar configuration is clearly untenable. We consider therefore the case where the cable configuration is not necessarily planar but is still time independent with respect to a reference frame translating with the constant velocity of the ship. In analogy with previous terminology, we call this the three-dimensional stationary model.

Assume there is a constant velocity ocean current in each of a finite number of layers. Let the vector $\vec{V}_w$ denote the ocean-current velocity in a reference layer. In the stationary situation the velocity of the cable

* Because of asymmetries caused by the helical armor wire or because of minor out-of-roundness, it is conceivable that a sidewise drag force might develop which would cause the cable to move out of the ship's velocity-gravity plane. For a report of experimental observations of such yawing in wire stranded cables, see Reference 11.

configuration is everywhere the velocity of the ship, which we denote by the vector $\vec{V}$. Hence the velocity $\vec{V}'$ of the water with respect to the cable configuration in the reference layer is

$$\vec{V}' = \vec{V}_w + (-\vec{V}) = \vec{V}_w - \vec{V}.$$

Further, in this layer we choose a set of coordinate axes ξ, η, ζ translating at the velocity $\vec{V}$ as follows: The ξ axis has the direction of $-\vec{V}'$, while η is measured vertically upward, and ζ is perpendicular to η and ξ so that the axes ξ, η, ζ form a right-handed system. A plan view of this configuration is shown in Fig. 26. We have denoted the angle between $\vec{V}$ and $\vec{V}_w$ by β , while the angle between the ξ axis and $\vec{V}$ is denoted by ϕ . (The distances d and e refer to a subsequent section.) To describe the cable configuration with respect to the ξ, η, ζ axes, we use the spherical polar coordinates θ and ψ shown in Fig. 27. (The $\vec{t}, \vec{u}, \vec{v}$ vectors are discussed in Appendix F.)

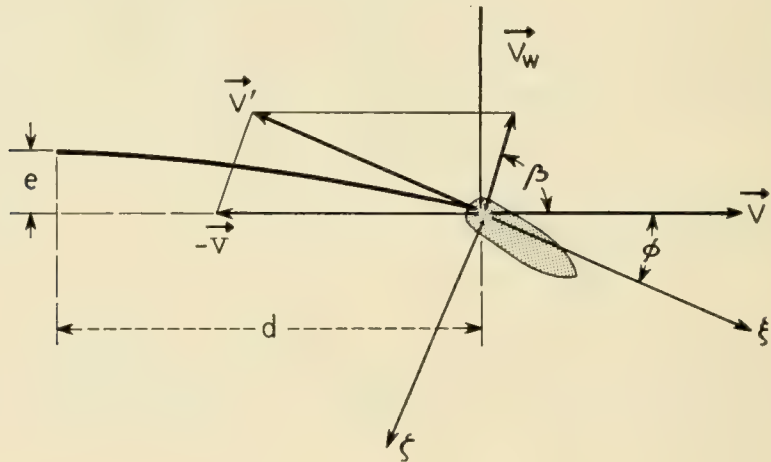


Fig. 26 — Plan view of the coordinate system for the three-dimensional stationary model.

As in the two-dimensional case, we resolve the velocity of the water with respect to a cable element in the reference layer into a component V_N normal to the cable and a component V_t tangential to the cable, and associate with V_N and V_t the drag forces D_N and D_T . The resulting differential equations, which are derived in detail in Appendix F, are the following:

$$(T - \rho_c V_c^2) \frac{d\theta}{ds} + w \Lambda' (\cos^2 \psi \sin^2 \theta + \sin^2 \psi)^{1/2} \cos \psi \sin \theta - w \cos \theta = 0, \quad (a)$$

$$(T - \rho_c V_c^2) \cos \theta \frac{d\psi}{ds} \quad (47)$$

$$+ w\Lambda'(\cos^2 \psi \sin^2 \theta + \sin^2 \psi)^{1/2} \sin \psi = 0, \quad (b)$$

$$\frac{dT}{ds} + D_T - w \sin \theta = 0, \quad (c)$$

where $\Lambda' = C_D \rho dV'^2/2$, and V' is the magnitude of $\vec{V}'$.

In addition, connecting the coordinates $\xi(s)$, $\eta(s)$, and $\zeta(s)$ of a point s along the cable with the angles θ and ψ we have the geometric relationships

$$\frac{d\xi(s)}{ds} = \cos \theta \cos \psi, \quad (a)$$

$$\frac{d\eta(s)}{ds} = \sin \theta, \quad (b) \quad (48)$$

$$\frac{d\zeta(s)}{ds} = -\cos \theta \sin \psi. \quad (c)$$

Two important general results follow from (47) and (48). For one, if the tangential drag force D_T is negligibly small, (48b) substituted into equation (47c) yields upon integration

$$T = T_0 + w\eta, \quad (49)$$

where T_0 is the tension at $\eta = 0$. Hence, if η is measured from the ocean

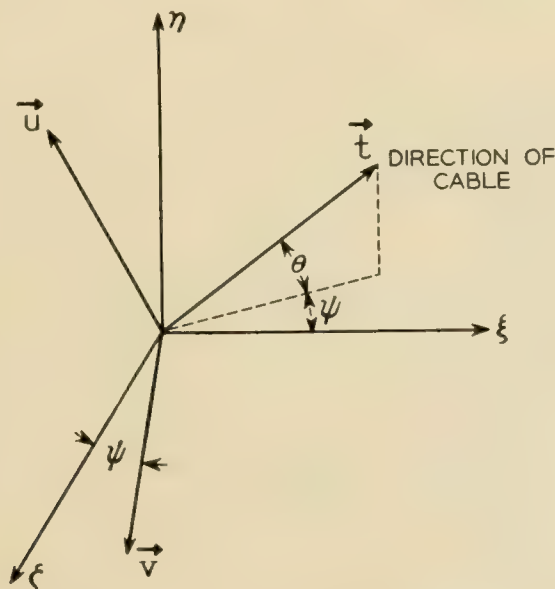


Fig. 27 — Definition of the spherical polar coordinates θ and ψ and the unit vectors $\vec{t}$, $\vec{u}$, and $\vec{v}$.

surface, and if at the bottom ($\eta = -h$) the tension is zero, the tension at the ship is essentially wh , regardless of the nature of the normal drag forces. Since in most laying situations for present cables, the tangential drag force can be reasonably neglected, this fact provides a convenient over-all check on the laying process. That is, if the cable is being laid with excess, the tension at the ship for *any* stationary cable configuration, *planar* or *non-planar*, should be essentially wh . Any marked increase of tension over the wh value necessarily means the bottom tension is non-zero and insufficient cable is being paid out.

The second important result is derived in Appendix F. This result is that if the bottom ocean layer in our model is devoid of cross currents, and if the bottom tension is zero, then, for the boundary conditions which are normally observed, the cable configuration in the bottom layer is a straight line. Further, this straight line is in the plane formed by the ship's velocity vector $\vec{V}$ and the gravity vector. Hence, for example, in laying with excess in a sea which contains surface currents, the cable configuration in the lower, current-free portion will be a straight line in a vertical plane parallel to the resultant velocity of the ship. The laid cable will be parallel to the ship's path, but displaced a certain distance from it. Thus, because the lower portion is a straight line, our previous results about the kinematics of straight-line laying still apply. Only they now are pertinent to the displaced bottom contour rather than to the contour which lies directly beneath the ship.

7.2 Perturbation Solution for a Uniform Cross-Current

Cross-currents are commonly confined to a region near the ocean surface. It is of interest therefore to determine for such surface currents the distance e (Fig. 26) which the laid cable will be displaced from the path of the ship. In Appendix F we consider the problem for a cross-current of uniform but comparatively small velocity. In addition, we determine the distance d (Fig. 26) back of the ship at which the cable leaves the upper, cross-current stratum and assumes the straight-line configuration it has in the lower stratum. Let us assume for the sake of reference that the resultant ship velocity V is due east, and that the cross-current V_u is inclined at an angle β to the north (Fig. 26). The resultant velocity V' of the water with respect to the cable in the surface stratum has the magnitude therefore of

$$V' = [(V - V_u \cos \beta)^2 + (V_u \sin \beta)^2]^{\frac{1}{2}}, \quad (50)$$

and is inclined at the angle φ from due west, where

$$\tan \varphi = \frac{V_w \sin \beta}{V - V_w \cos \beta}. \quad (51)$$

Associated with V' we have a critical angle α' which is given approximately by H/V' or exactly by (10) or (11).

In terms of φ , α' , and α the analysis of Appendix F yields the following values of d and e .

$$\begin{aligned} d &= h' \left(\text{ctn} \alpha' - \frac{\Delta \alpha}{2 \cos^2 \alpha'} \frac{h - h'}{h'} \left[1 - \left(1 - \frac{h'}{h} \right)^{2 \text{ctn}^2 \alpha'} \right] \right) \quad (a) \\ e &= h' \varphi \text{ctn} \alpha' \left(1 - \frac{h - h'}{h'} \tan^2 \alpha' \left[1 - \left(1 - \frac{h'}{h} \right)^{\text{ctn}^2 \alpha'} \right] \right), \quad (b) \end{aligned} \quad (52)$$

where $\Delta \alpha = \alpha - \alpha'$ is the difference of lower and upper stratum critical angles, h' is the depth of the upper, cross-current stratum and h is the total depth.* Curves from which d and e may be evaluated are given in Figs. 28 and 29 in dimensionless form. To illustrate their application we consider the following.

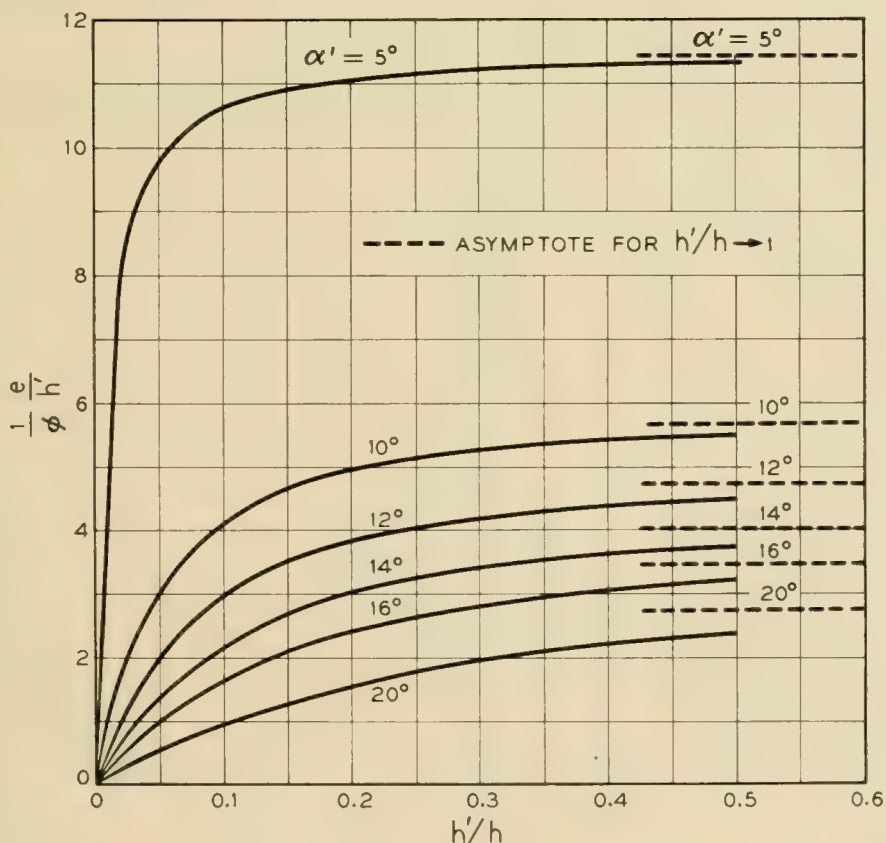


Fig. 28 — Distance the laid cable is offset from the ship's path.

* In Appendix F the equation of the space curve formed by the cable in the upper stratum is also given.

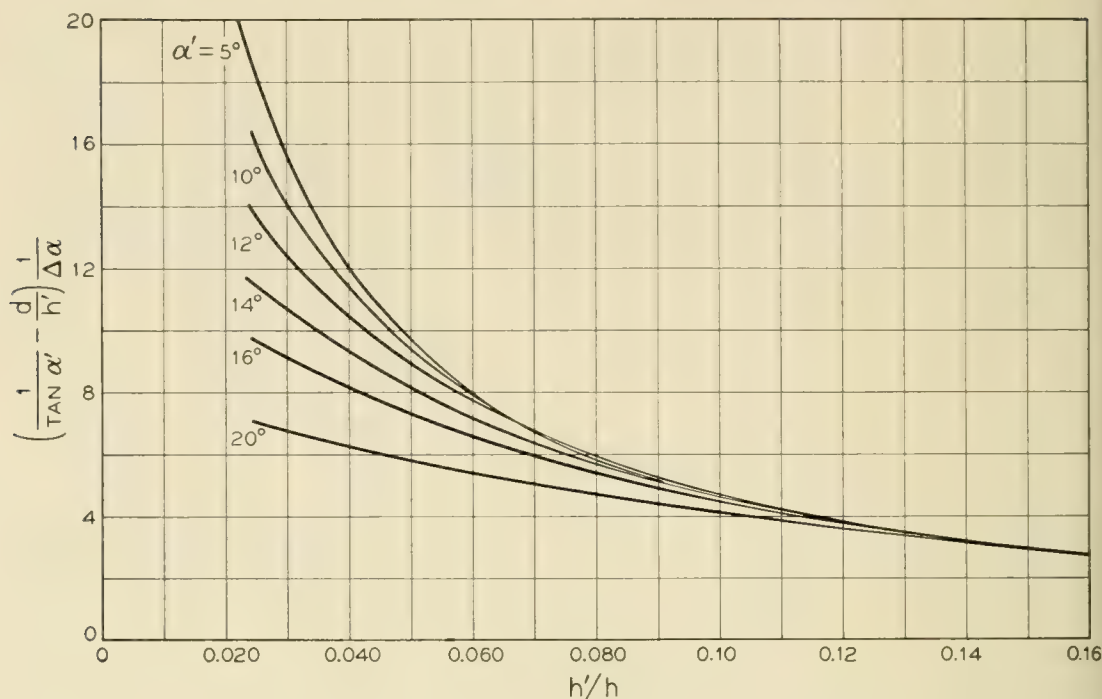


Fig. 29 — Distance behind the ship at which the cable enters the lower stratum.

Example: A ship is laying cable No. 2 at a depth of 6,000 ft at a resultant ground speed of 6 knots due east. There is a one knot cross-current 600 feet deep running 30° east of north. Find e and d .

Here $\beta = 60^\circ$ and we obtain from (50) by a simple computation $V' = 5.57$ knots. Also since for cable No. 2 $H = 70$ degree-knots,

$$h'/h = 600/6000 = 0.1,$$

$$\alpha' = 70/5.57 = 12.3 \text{ degrees.}$$

By interpolation, we find from Figs. 28 and 29

$$\frac{1}{\varphi} \frac{e}{h'} = 2.7,$$

$$\left(\frac{1}{\tan \alpha'} - \frac{d}{h'} \right) \frac{1}{\Delta \alpha} = 4.7.$$

Equation (51) yields in turn $\varphi = 0.156$ radians, and $\Delta \alpha$ is given by

$$\Delta \alpha = \frac{70}{V} - \frac{70}{V'} = -0.6 \text{ degrees} = -0.010 \text{ radians.}$$

Hence, we have

$$e = 3.6 h' \varphi = 2.7 \times 600 \times 0.156 = 253 \text{ ft,}$$

$$d = h' \left[\frac{1}{\tan \alpha'} - 4.7 \Delta \alpha \right] = 6000 [4.59 + 4.7 \times 0.010] = 2800 \text{ ft.}$$

APPENDIX A

Discussion of the Two-Dimensional Stationary Configuration for Zero Bottom Tension

We assume again that the tangential drag D_T depends only on the relative tangential velocity V_T , and we consider in a T, θ plane the solution trajectories of equations (18). These trajectories satisfy the equation

$$\frac{dT}{d\theta} = \frac{(\sin \theta - D_T/w)}{\cos \theta - \Lambda \sin \theta |\sin \theta|} (T - \rho_c V_c^2), \quad (53)$$

and are periodic in θ with a period of 2π . In Fig. 30 we have plotted the solution trajectories qualitatively for $(\alpha - \pi) \leq \theta \leq (\alpha + \pi)$. It is seen that the trajectories are either the vertical straight lines $\theta = \alpha, \theta = \alpha \pm \pi$ or they lie completely within one of four regions, labelled I, II, III, or IV, which are bounded by these vertical lines and the horizontal line $T = \rho_c V_c^2$. The trajectory $\theta = \alpha$ corresponds to the straight-line laying configuration, while the trajectories $\theta = \alpha \pm \pi$ correspond to Shea's straight-line recovery method.

Examine now the trajectories in Regions II and III at a point of which $T = 0$. As J. F. Shea has pointed out, these trajectories all lie below the line $T = \rho_c V_c^2$. On the other hand, the trajectory $\theta = \alpha$ con-

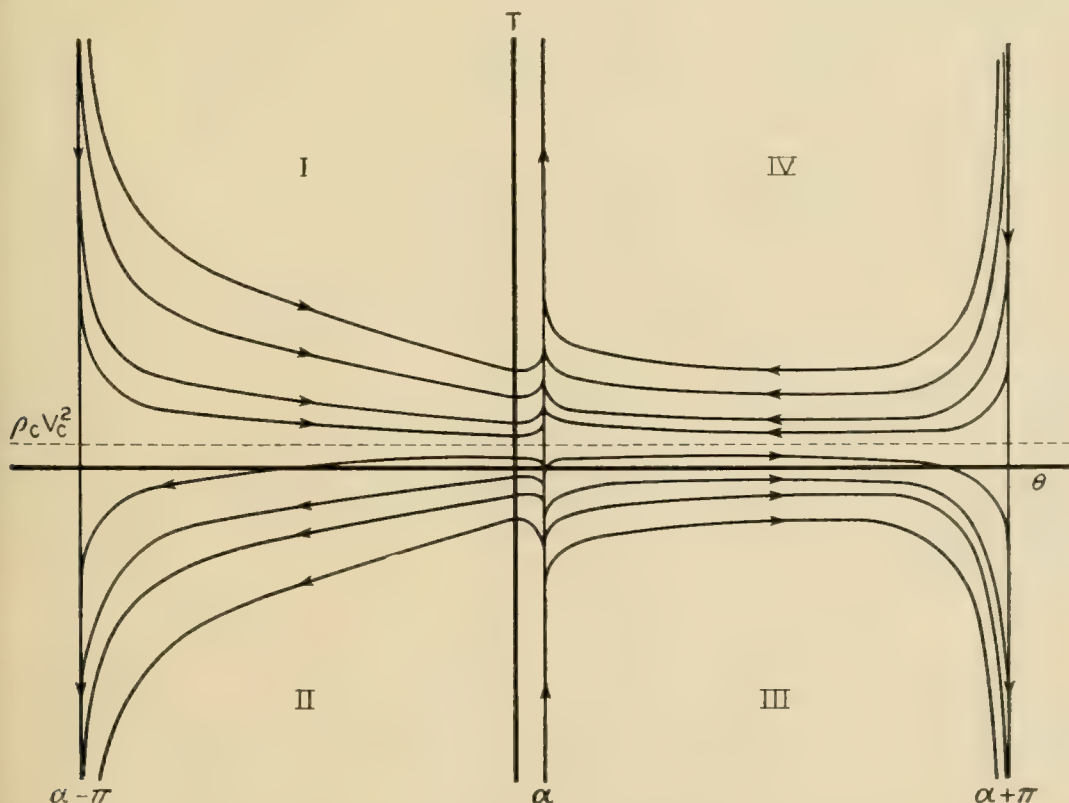


Fig. 30 — Qualitative representation of the solution trajectories of the two-dimensional stationary model.

tains all values of T . Hence, according to the stationary model, the only cable configuration for laying which has the value $T = 0$ and values of $T \geq \rho_c V_c^2$ is the straight line inclined at the critical angle α . The magnitude of $\rho_c V_c^2$ is small. For example, for cable No. 2 paid out at six knots $\rho_c V_c^2$ is roughly six pounds, and for conditions approximating stationary laying the observed tensions at the ship are in practice always many times the $\rho_c V_c^2$ value. For such magnitudes of shipboard tensions and a zero bottom tension, the two-dimensional stationary model thus yields the straight line as the only possible cable configuration.

However, the empirical fact that $T > \rho_c V_c^2$ does not guarantee that the shipboard tension must be greater than $\rho_c V_c^2$. We might somehow contrive to lay at a zero bottom tension with $T < \rho_c V_c^2$ and with the cable in one of the non-straight line configurations of Regions II or III.

Consider the cable configuration lying in Region II. From Fig. 7 it can be seen that the vertical velocity of a cable element is given by

$$\frac{dy}{dt} = -V_{\text{vert}} = -V_c \sin \theta,$$

where y is measured upward. Hence, of the possible trajectories for which the bottom tension is zero only those for which the bottom cable angle θ_0 is between zero and π correspond to cable laying. For region II, therefore we need consider only the trajectories in the range $0 \leq \theta_0 < \alpha$ at $T_0 = 0$. From (20c) the maximum value of y_m for these trajectories is given by

$$y_m = \frac{\rho_c V_c^2}{w} \int_0^{\theta_0} \left\{ \frac{\sin \xi}{(\cos \xi - \Lambda \sin^2 \xi)} \right. \\ \left. \times \exp \left[- \int_{\xi}^{\theta_0} \frac{w \sin \eta - D_T(\eta)}{w (\cos \eta - \Lambda \sin^2 \eta)} d\eta \right] \right\} d\xi. \quad (54)$$

Let $(D_T)_m$ be the maximum value of D_T , $0 \leq \eta < \alpha$. With D_T set equal to $(D_T)_m$, the right-hand side of (54) gives an upper bound on y_m . This substitution further allows one to evaluate the right-hand side of this equation in terms of standard integrals. The result yields the following upper bound on y_m :

$$y_m < 2.1 \frac{\rho_c V_c^2}{w} \frac{1}{1 - r},$$

where

$$r = \frac{(D_T)_m}{w \sin \alpha}.$$

In general, this upper bound will be much less than the laying depth. For example, for cable No. 2 being laid with 6 per cent slack at 6 knots $y_m < 12.5$ feet. That is, the cable configurations corresponding to Region II do not reach the ocean surface. Hence these solutions of the stationary model do not in general satisfy all the required boundary conditions and can be discarded.

Similarly, in Region III, the laying trajectories for which $T_0 = 0$ are in the range $\alpha < \theta_0 \leq \pi$. Consider those for which $\theta_0 < \pi/2$. We get for these trajectories

$$y_{\pi/2} = \frac{\rho_c V_c^2}{w} \int_{\theta_0}^{\pi/2} \left\{ \frac{\sin \xi}{(\Lambda \sin^2 \xi - \cos \xi)} \right. \\ \left. \times \exp \left[- \int_{\theta_0}^{\xi} \frac{w \sin \eta - D_T(\eta)}{w (\Lambda \sin^2 \eta - \cos \eta)} d\eta \right] \right\} d\xi, \quad (55)$$

where $y_{\pi/2}$ is the value of y at $\theta = \pi/2$. Let m be the minimum value of $\sin \eta - (D_T(\eta)/w)$ in the range $\alpha < \eta \leq \pi/2$. If, as in the usual case, m is positive, we can obtain an upper bound on $y_{\pi/2}$ by replacing $\sin \eta - (D_T(\eta)/w)$ by m in the right-hand side of (55). By this means we find that

$$y_{\pi/2} < \frac{\rho_c V_c^2}{w} \frac{2(1 + \cos^2 \alpha)}{m \tan \alpha/2}.$$

For cable No. 2 being laid with 6 per cent slack at 6 knots this relation yields $y_{\pi/2} < 1,100$ feet. So in the usual laying depths, which are many times greater than $y_{\pi/2}$, the configurations in Region III for which $T_0 = 0$ correspond to a value of θ at the surface greater than $\pi/2$, or to cable being paid out *in front* of the ship during laying. It is doubtful whether such configurations would be stable and, at any rate, doubtful whether cable would ever be laid in such a manner. Hence, we conclude that these $T_0 = 0$ solutions of Regions II and III will in general be mathematical curiosities, and that the only realistic laying solution of the stationary model for which the bottom tension is zero is the straight line $\theta \equiv \alpha$.

APPENDIX B

Computation of the Transverse Drag Coefficient and the Hydrodynamic Constant of a Smooth Cable from Published Data

From (6) of Section 3.2 we obtain the relationship

$$\frac{V \sin \alpha}{\sqrt{\cos \alpha}} = \left(\frac{2w}{C_D \rho d} \right)^{1/2} = H, \quad (56)$$

where H is the hydrodynamic constant. Also, we define the Reynolds number for flow normal to the cable in the usual way;

$$N_R = \frac{dV \sin \alpha}{\nu}, \quad (57)$$

ν being the kinematic viscosity of water. For smooth cable we now assume that the drag coefficient C_D depends on N_R in the same way as for flow around a smooth cylinder of infinite length. Experimental data for this relationship, namely,

$$C_D = C_D(N_R) \quad (58)$$

are available in the literature and have been collated by Eisner.¹⁵

For a given velocity V , (56) through (58) represent three equations for the unknowns α , C_D , and N_R . In general, the solution of these equations depends on V , thus contradicting the assumption made in Section 3.2 that C_D , and therefore H , are constants independent of V . However, for sufficiently large V we can expect the resulting α to be small so that $\sqrt{\cos \alpha} \approx 1$. In this case (56) and (57) combine to give

$$C_D = \frac{2wd}{\rho \nu^2} \frac{1}{N_R^2}, \quad (59)$$

which together with (58) yields two equations for C_D and N_R that are indeed independent of V . In laying, V will normally be large enough for this approximation to hold.

Equations (58) and (59) are in turn easily solved graphically by finding the intersection on log-log paper of the curves, C_D versus N_R , that these equations represent. Since ρ and ν are properties only of the water, we see that C_D is a function only of the product wd of the unit weight of the cable times its diameter. In Fig. 31 we have plotted the resulting values of C_D for wd ranging from 10^{-7} to 10 pounds. For this computation we have assumed sea water at 32°F with an assumed density of 1.994 slugs per cubic foot and a kinematic viscosity of 2.006×10^{-5} ft²/sec.

For other than large values of V , (56) through (58) can be readily solved if one interchanges the roles of α and V , that is, if one considers α as given and V as unknown. Equations (56) and (57) can again be combined in this case to give

$$C_D = \frac{2wd \cos \alpha}{\rho v^2} \frac{1}{N_R^2}, \tag{60}$$

and with $wd \cos \alpha$ a known number, (60) and (58) can be solved for C_D and N_R as before. Thus, one can obtain C_D from Figure 31 by merely reading $wd \cos \alpha$ rather than wd on the abscissa. Knowing C_D one can solve for V from (56).

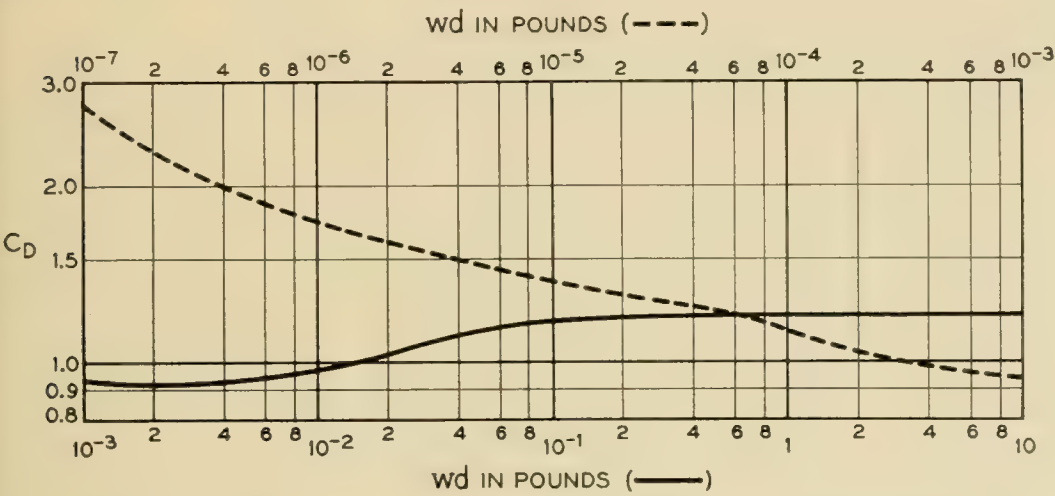


Fig. 31 — Variation of C_D with wd for cables of smooth exterior.

The results of such a computation are shown in Table II for cable No. 1. For $V > 1.5$ knots the experimentally determined H is 64.0 degree-knots. The corresponding computed values of H , ranging from 67.4 to 70.0 degree-knots, compare favorably with this experimental value. Over the entire range of V , from 0.25 to 10.00 knots, the variation in H is about 4 per cent. This small variation makes the assumption of Section 3.2 that H is a constant for all V appear reasonable, especially since we can only hope to use this computation in a preliminary design before the hydrodynamic properties of a cable are established by experiment.

TABLE II — COMPUTED VALUES OF N_R , C_D , AND H FOR CABLE NO. 1

V (knots)	N_R	C_D	H (deg.-knots)
0.25	1.25×10^3	0.935	67.5
0.50	2.50	0.922	70.0
1.50	5.20	0.965	67.8
3.00	5.80	0.985	67.5
10.00	6.05	1.000	67.4

In Table III we have indicated computed high-velocity H values as a function of the unit weight w and the diameter d of a cable. Table IV shows the computed high-velocity C_D and H for various gauge piano wire. The American Steel and Wire gauge scale is used in this tabulation.

TABLE III — COMPUTED VALUES OF THE HYDRODYNAMIC CONSTANT H IN DEGREE-KNOTS FOR SMOOTH CABLE

Submerged Weight in lb/ft	Diameter in Inches									
	0.5	0.75	1.00	1.25	1.50	1.75	2.0	2.5	3.0	4.0
0.1	54.5	44.2	37.9	33.7	30.6	28.1	26.2	23.2	21.0	17.9
0.2	75.9	61.1	52.3	46.4	41.9	38.5	35.8	31.6	28.5	24.4
0.3	91.7	73.6	62.9	55.6	50.2	46.1	42.8	38.0	34.4	29.7
0.4	104.7	83.9	71.5	63.1	57.1	52.5	48.9	43.5	39.6	34.3
0.5	115.9	92.6	78.9	69.8	63.3	58.3	54.4	48.5	44.3	38.3
0.6	125.8	100.4	85.6	75.9	68.9	63.6	59.5	53.2	48.5	42.0
0.7		107.6	91.9	81.6	74.2	68.7	64.2	57.4	52.4	45.3
0.8		114.2	97.8	87.0	79.3	73.4	68.6	61.3	55.9	48.4
0.9		120.6	103.4	92.1	84.1	77.8	72.7	65.0	59.3	51.3
1.0		126.5	108.7	97.1	88.6	82.0	76.6	68.5	62.5	54.1
2.0			153.3	137.0	125.0	115.7	108.2	96.7	88.2	76.3
3.0				167.6	152.9	141.5	132.3	118.2	107.9	93.3

TABLE IV — COMPUTED C_D AND H VALUES OF PIANO WIRE

Gauge (Am. Steel & Wire)	Dia. (inches)	C_D	H (deg- knots)
0	0.009	2.49	10.7
5	0.014	1.91	15.2
10	0.024	1.56	22.1
15	0.035	1.39	28.2
20	0.045	1.31	33.0
25	0.059	1.24	38.7
30	0.080	1.14	47.2
35	0.106	1.02	57.1
40	0.138	0.970	67.1

APPENDIX C

Some Approximate Solutions for Laying and Recovery

c.1 Laying

We assume that the tangential drag and the centrifugal forces are negligible. Then, since for laying $0 \leq \theta \leq \pi$, (18a) by virtue of (21) becomes

$$\left(\frac{T_0}{w} + y\right) \frac{d\theta}{ds} + \Lambda \sin^2 \theta - \cos \theta = 0. \tag{61}$$

Let the origin of an x, y coordinate system be at the cable touchdown point (Fig. 8). Further, let x be the x coordinate of a point along the cable configuration and s the corresponding distance along the cable from the origin. If we define

$$\Delta = s - x, \quad (62)$$

then

$$\frac{d\Delta}{dy} = \frac{ds}{dy} - \frac{dx}{dy} = \tan \frac{\theta}{2}, \quad (63)$$

and

$$\frac{dy}{ds} = \sin \theta. \quad (64)$$

By means of (63) and (64), (61) transforms to

$$(\bar{T}_0 + \bar{y}) \frac{d\bar{\Delta}}{d\bar{y}} \frac{d^2\bar{\Delta}}{d\bar{y}^2} + \frac{1}{4} \frac{(d\bar{\Delta})^4}{(d\bar{y})} + \Lambda \frac{(d\bar{\Delta})^2}{(d\bar{y})} - \frac{1}{4} = 0, \quad (65)$$

where we have in addition introduced the non-dimensional variables

$$\bar{T}_0 = T_0/wh,$$

$$\bar{\Delta} = \Delta/h,$$

$$\bar{y} = y/h.$$

Here h is the ocean depth at the touchdown point. Using the condition that $\theta = 0$ at $y = 0$, which implies $d\Delta/dy = 0$ at $y = 0$, we get upon integrating (65)

$$\frac{d\bar{\Delta}}{d\bar{y}} = \tan \frac{\alpha}{2} \left(\frac{1 - [\bar{T}_0/(\bar{T}_0 + \bar{y})]^\gamma}{1 + [\bar{T}_0/(\bar{T}_0 + \bar{y})]^\gamma \tan^4 \frac{\alpha}{2}} \right)^{1/2}, \quad (66)$$

where

$$\gamma = \frac{2 - \sin^2 \alpha}{\sin^2 \alpha}. \quad (67)$$

The usual range of the critical angle α is between 10 and 30 degrees. Also

$$0 \leq \frac{\bar{T}_0}{\bar{T}_0 + \bar{y}} \leq 1.$$

Therefore, we approximate the denominator of equation (67) by unity.

With the boundary condition that $\Delta = s - x = 0$ at $y = 0$, we thus obtain

$$\bar{\Delta}(1) = \bar{S} - \bar{X} = \tan \frac{\alpha}{2} \int_0^1 (1 - [\bar{T}_0/(\bar{T}_0 + \xi)]^\gamma)^{\frac{1}{2}} d\xi, \quad (68)$$

where $\bar{S}$ and $\bar{X}$ are the dimensionless values of s and x at the ship.

Next we let

$$\omega = s + x, \quad \bar{\omega} = \omega/h.$$

Then we have

$$\frac{d\bar{\omega}}{d\bar{y}} = 1 \bigg/ \frac{d\bar{\Delta}}{d\bar{y}},$$

and, as can be seen from (66) through (68),

$$\bar{\omega}(1) = \bar{S} + \bar{X} = \operatorname{ctn} \frac{\alpha}{2} \int_0^1 (1 - [\bar{T}_0/(\bar{T}_0 + \xi)]^\gamma)^{-\frac{1}{2}} d\xi. \quad (69)$$

For convenience we define u and R by

$$u = \frac{\bar{T}_0}{\bar{T}_0 + \xi},$$

$$R = \frac{\bar{T}_0}{1 + \bar{T}_0}.$$

In terms of u and R (68) and (69) become

$$\bar{\Delta}(1) = \bar{T}_0 \tan \frac{\alpha}{2} \int_R^1 \frac{(1 - u^\gamma)^{\frac{1}{2}}}{u^2} du, \quad (70)$$

$$\bar{\omega}(1) = \bar{T}_0 \operatorname{ctn} \frac{\alpha}{2} \int_R^1 \frac{du}{u^2(1 - u^\gamma)^{\frac{1}{2}}}. \quad (71)$$

Further, integration by parts gives

$$\int_R^1 \frac{(1 - u^\gamma)^{\frac{1}{2}}}{u^2} du = \frac{(1 - R^\gamma)^{\frac{1}{2}}}{R} + \frac{\gamma}{2} \int_R^1 \frac{(1 - u^\gamma)^{\frac{1}{2}}}{u^2} du - \frac{\gamma}{2} \int_R^1 \frac{du}{u^2(1 - u^\gamma)^{\frac{1}{2}}}.$$

Combining the above three equations and making the approximation $(1 - R^\gamma)^{\frac{1}{2}} \approx 1$, we find

$$\left(1 - \frac{\gamma}{2}\right) \bar{\Delta}(1) + \frac{\gamma}{2} \tan^2 \frac{\alpha}{2} \bar{\omega}(1) = (1 + \bar{T}_0) \tan \frac{\alpha}{2}. \quad (72)$$

Thus $\bar{\Delta}(1)$ and $\bar{\omega}(1)$ are related, and we need evaluate only one of the quantities numerically by means of equation (70) or (71) in order to compute both $\bar{\Delta}(1)$ and $\bar{\omega}(1)$, and hence $\bar{S}$ and $\bar{X}$.

The singularity at $u = 1$ makes the numerical evaluation of the integral in (71) cumbersome. Therefore we consider the evaluation of $\bar{\Delta}(1)$. But for convenience of numerical calculation we write (70) as

$$\bar{\Delta}(1) = -\bar{T}_0 \tan \frac{\alpha}{2} \int_R^1 (1 - u^\gamma)^{\frac{1}{2}} d\left(\frac{1 - u}{u}\right),$$

and integrate by parts to get

$$\bar{\Delta}(1) = \tan \frac{\alpha}{2} \left((1 - R^\gamma)^{\frac{1}{2}} - \frac{\gamma}{2} \bar{T}_0 \int_R^1 \frac{1 - u}{u^2} \frac{u^\gamma}{(1 - u^\gamma)^{\frac{1}{2}}} du \right).$$

We note again that $(1 - R^\gamma)^{\frac{1}{2}} \approx 1$. Further, essentially all of contribution to the integral in this equation occurs near $u = 1$ because of the large value of γ . On the other hand, the values of $\bar{T}_0$ which are of interest will normally be smaller than unity. Hence R , the lower limit, will normally be less than one-half, and thus will be outside of the region of significant contribution to the integral. Therefore, we can take the integral to be a constant for a given α . Denoting this integral by n and combining these considerations we obtain

$$\bar{\Delta}(1) = \tan \frac{\alpha}{2} - \frac{\gamma}{2} n \tan \frac{\alpha}{2} \bar{T}_0. \quad (73)$$

Finally solving for $\bar{S}$ and $\bar{X}$ from (68) and (69), we find

$$\bar{S} = \frac{1}{\sin \alpha} + \kappa \bar{T}_0,$$

$$\bar{X} = \frac{1}{\tan \alpha} + \lambda \bar{T}_0,$$

which are a dimensionless form of (24). For brevity we have written

$$\kappa = \frac{1}{2} \left(\frac{2}{\lambda} \left[1 - \frac{\gamma}{2} \left(\frac{\gamma}{2} - 1 \right) n \right] \operatorname{ctn} \frac{\alpha}{2} - \frac{\gamma}{2} n \tan \frac{\alpha}{2} \right),$$

$$\lambda = \frac{1}{2} \left(\frac{2}{\gamma} \left[1 - \frac{\gamma}{2} \left(\frac{\gamma}{2} - 1 \right) n \right] \operatorname{ctn} \frac{\alpha}{2} + \frac{\gamma}{2} n \tan \frac{\alpha}{2} \right).$$

Since the integral n and the constant γ depend only on α , the constants κ and λ are also functions of α only. We have evaluated n by numerical integration and have plotted the resulting values of κ and $\lambda - \kappa$ in Fig. 10.

c.2 Recovery

In the conventional recovery situation we have as boundary conditions

$$\begin{aligned}\theta &= 0 & \text{at } y &= 0, \\ \theta &= -\alpha_s & \text{at } y &= h,\end{aligned}\tag{74}$$

where α_s is the incidence angle of the cable at the ship (Fig. 8). With these boundary conditions, the development leading to (66) yields (26) for the relationship between shipboard tension $\bar{T}_s$ and the incidence angle α_s .

To evaluate S and X , the values of s and x at the ship, we use (19), (20a) and (20b). Again we simplify by assuming $D_T = \rho_c V_c^2 = 0$. Further, since for recovery $-\pi \leq \theta \leq 0$, we may write $|\sin \xi| = -\sin \xi$. This gives

$$\begin{aligned}\bar{S} &= \bar{T}_0 \int_0^{-\alpha_s} \left[\frac{1}{\cos \xi + \Lambda \sin^2 \xi} \exp \int_0^\xi \frac{\sin \eta}{\cos \eta + \Lambda \sin^2 \eta} d\eta \right] d\xi, \\ \bar{X} &= \bar{T}_0 \int_0^{-\alpha_s} \left[\frac{\cos \xi}{\cos \xi + \Lambda \sin^2 \xi} \exp \int_0^\xi \frac{\sin \eta}{\cos \eta + \Lambda \sin^2 \eta} d\eta \right] d\xi.\end{aligned}\tag{75}$$

The dimensionless bottom tension $\bar{T}_0$ is computed from (26). The integrals appearing in (75) have been evaluated numerically. The results are shown in Figs. 14 and 15.

APPENDIX D

Analysis of the Effect of Ship Motion

D.1 Formulation of the Differential Equations

To analyze the effect of ship motion on cable tensions, we use the model shown in Fig. 32. We assume the cable is a perfectly flexible and elastic string whose motion is planar. The distance L along the cable from the ship to the point of entry into the water is taken as constant, and the longitudinal damping as negligible.

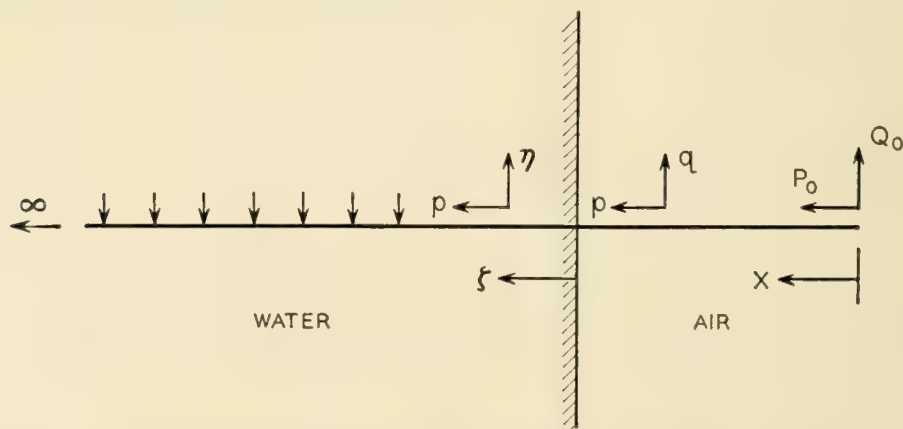


Fig. 32 — Model used for the analysis of ship motion tensions.

Unlike the solution of the basic stationary model, the complete solution of this model is not simple. To make the problem tractable, we shall make further simplifying assumptions. Although these assumptions may seem reasonable, they must be ultimately justified by comparison of experience with predicted results.

Force diagrams of a differential element of cable are shown in Figs. 33(a) and (b) for the two regions, air and water respectively. The notation is

- p = longitudinal displacement of a point of the cable,
- q = transverse displacement of a point of the cable (in air),
- η = transverse displacement of a point of the cable (in water),
- θ = the stationary angle, i.e. the angle the cable configuration makes with the ship velocity in the absence of ship motion,
- φ = deviation from the stationary angle – positive in the clockwise direction,
- s = distance along the stretched cable,
- x = distance along the unstretched cable (in air),
- ζ = distance along the unstretched cable (in water),
- w_a = weight per unit length of cable in air.

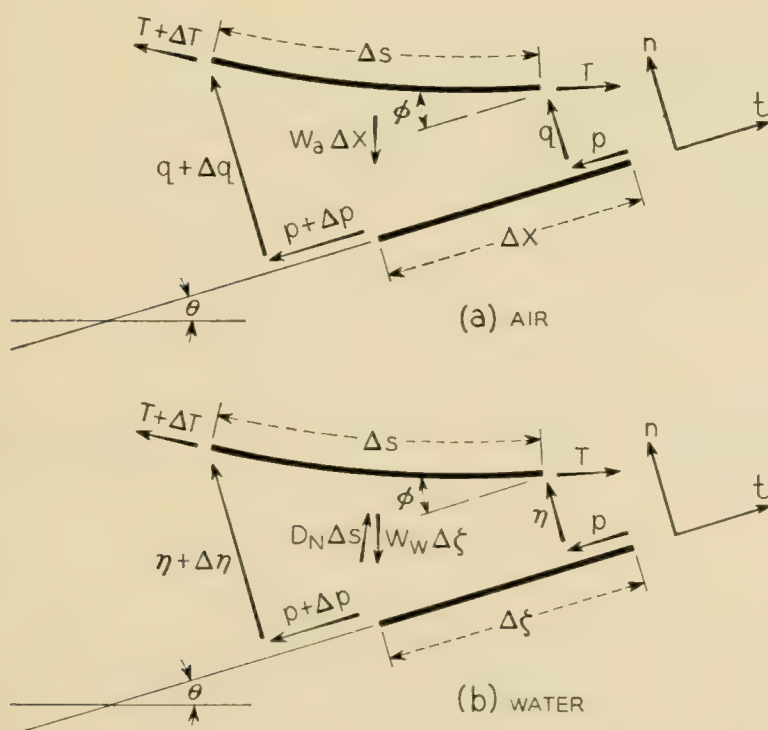


Fig. 33 — Diagram of forces acting on a cable element in air and in water.

Summing forces along the directions t (tangential) and n (normal) shown in Fig. 33, dividing by Δx (air) or $\Delta \zeta$ (water) and letting $\Delta x \rightarrow 0$ and $\Delta \zeta \rightarrow 0$, we obtain the following equations of equilibrium*

Air:

$$\begin{aligned} T\varphi_x - w_a \cos(\theta - \varphi) &= \rho_c(q_{tt} \cos \varphi - p_{tt} \sin \varphi), \\ T_x + w_a \sin(\theta - \varphi) &= \rho_c(q_{tt} \sin \varphi + p_{tt} \cos \varphi), \end{aligned} \quad (76)$$

Water:

$$\begin{aligned} T\varphi_\zeta + D_N s_\zeta - w \cos(\theta - \varphi) &= \rho_w(q_{tt} \cos \varphi - p_{tt} \sin \varphi), \\ T_\zeta + w \sin(\theta - \varphi) &= \rho_w(q_{tt} \sin \varphi + p_{tt} \cos \varphi). \end{aligned} \quad (77)$$

Here, ρ_c denotes the mass per unit length of the cable in air. As is known from hydrodynamic theory, in order to accelerate a body through a fluid, one must change not only the momentum of the body but that of some of the surrounding fluid as well. Thus the body has a virtual or apparent mass in addition to its intrinsic mass. In the first (77), the equation of equilibrium in the normal direction in water, we accordingly use ρ_w , given by

$$\rho_w = \rho_c + \frac{\pi}{4} d^2 \rho$$

as the intrinsic plus virtual mass per unit length of cable moving through water. The quantities d and ρ are the outer diameter of the cable and mass density of the water, respectively, and the quantity $(\pi/4)d^2\rho$ is the virtual mass of a unit length circular cylinder moving transversely through water.

We take for the normal drag force per unit length

$$D_N = \frac{C_D \rho d}{2} V_N |V_N|. \quad (78)$$

Here V_N is the normal component of velocity of the water relative to the cable, i.e.,

$$V_N = V \sin(\theta - \varphi) + u_t \sin \varphi - \eta_t \cos \varphi, \quad (79)$$

and $C_D \rho d/2$ is a constant.

The quantities s and φ are given by the following geometric relations which can be obtained from Fig. 33:

* We use the subscript notation for differentiation throughout this section, e.g., $\varphi_x = \partial \varphi / \partial x$.

$$s_{\zeta} = [(1 + p_{\zeta})^2 + \eta_{\zeta}^2]^{\frac{1}{2}}, \quad (80)$$

$$\tan \varphi = \eta_{\zeta}/(1 + \eta_{\zeta}), \quad (81)$$

with similar expressions obtaining for the cable in air.

The tension T in turn is given by the Hooke's Law or stress-strain relation

$$\begin{aligned} T &= \overline{EA} \{[(1 + p_x)^2 + q_x^2]^{\frac{1}{2}} - 1\}, & (\text{air}) \\ T &= \overline{EA} \{[(1 + p_{\zeta})^2 + \eta_{\zeta}^2]^{\frac{1}{2}} - 1\}. & (\text{water}) \end{aligned} \quad (82)$$

As we indicated in Section 4.1, we shall assume that the extensile rigidities $\overline{EA}$ corresponding to complete restraint and no restraint to twisting will give the limiting values of the ship motion tension.

Equations (76) through (82) form a complete system in terms of the independent variables x or ζ and t . Formulating boundary conditions in terms of the coordinate x (or ζ) is awkward. This coordinate is measured along the unstretched cable so that a disturbance applied at the ship is applied at different x 's as the cable is paid out. At the same time, if the velocity of the pay-out is small compared to the significant wave velocity of the cable then we can plausibly neglect the paying out effect. As will be shown subsequently, in the problem at hand there are two significant wave velocities, roughly corresponding to transverse and longitudinal motion. The first of these is of the order of 200 ft/sec, while the second is of the order of 5,000–10,000 ft/sec. On the other hand, the pay-out velocity is of the order of 10 ft/sec. Hence, we take the pay-out velocity to be zero. This allows us to use (76) through (82) without further transformations and to identify x and ζ as coordinates fixed in the translating reference frame.

D.2 *Perturbation Equations*

We assume that the motion is a small perturbation about the undisturbed configuration of our model. To determine which terms of the differential equations are important in this case, we adopt the following procedure. Let

$$M = \max[(P_0 + P_1)^2 + Q_0^2]^{\frac{1}{2}},$$

where P_0 and Q_0 are displacements of the cable at the ship, and P_1 is the variation of the pay-out displacement from the mean. The quantity $e = M/L$ will normally be less than unity, and for no ship motion will be

zero. We write

$$P_0 + P_1 = af(t),$$

$$Q_0 = bg(t),$$

where $f(t)$, $g(t)$ are some bounded functions of time and a and b are constants. We assume that for given $f(t)$ and $g(t)$, T , p , and q vary analytically with e , namely

$$\begin{aligned} T &= T_0 + eT_1 + e^2T_2 + \cdots, \\ q &= eq_1 + e^2q_2 + \cdots, \\ p &= p_0 + ep_1 + e^2p_2 + \cdots, \end{aligned} \quad (83)$$

with counterparts for the submerged cable. The stationary transverse deflection is further assumed zero, and therefore the series for q contains no q_0 term. Substituting, for example, (83) into (82) for air and equating like powers of e , we find

$$T_0 = \overline{EA}p_{0x}, \quad (a)$$

$$T_1 = \overline{EA}p_{1x}, \quad (b) \quad (84)$$

$$T_2 = \overline{EA} \left[p_{2x} + \frac{q_{1x}^2}{1 + p_{0x}} \right]. \quad (c)$$

Equation (84a) of this sequence shows that only longitudinal displacements are associated with stationary tensions, while (84b) indicates that for small ship motions cable tensions are independent of the transverse component of ship motion. To compute the effect of transverse motion, (84c) shows that terms of the order e^2 in p and e in q must be considered. We assume further that $1 + p_{0x} \approx 1$, since p_{0x} is the order of magnitude of a strain.

Equations (83) and (84) substituted into (76) yield with this approximation

$$w_a \cos \alpha = 0, \quad (a) \quad (85)$$

$$\overline{EA} p_{0xx} - \rho_a p_{0tt} + w_a \sin \alpha = 0, \quad (b)$$

$$q_{1rx} - \frac{1}{c_2^2} q_{1tt} = 0, \quad (a)$$

$$p_{1xx} - \frac{1}{c_1^2} p_{1tt} = 0, \quad (b) \quad (86)$$

$$p_{2rx} - \frac{1}{c_1^2} p_{2tt} = \frac{1}{c_1^2} q_{1tt}q_{1rx} - q_{1x}q_{1xx}, \quad (c)$$

where

$$c_1^2 = \overline{EA}/\rho_a,$$

$$c_2^2 = T_0/\rho_a.$$

For non-zero w_a and $\alpha \neq \pi/2$, (85a) cannot be satisfied. This is a consequence of the assumption of $q_0 = 0$. With $p_{0tt} = 0$, equation (b) implies in turn that $T_0 = \text{constant}$, which agrees with our model. For the submerged part of the cable, the equations do not yield a constant T_0 and thus contradict the assumed model. However, on the assumption that the transverse motion is confined to a region near the surface, we consider T_0 to be constant in the submerged part of the cable as well. We thus arrive at

$$\eta_{1\xi\xi} - \delta\eta_{1\xi} - \gamma\eta_{1t} - \frac{1}{\bar{c}_2^2}\eta_{1tt} = 0 \quad (\text{a})$$

$$p_{1\xi\xi} - \frac{1}{c_1^2}p_{1tt} = 0, \quad (\text{b}) \quad (87)$$

$$p_{2\xi\xi} - \frac{1}{c_1^2}p_{2tt} = \frac{1}{c_1^2}\eta_{1tt}\eta_{1\xi} - \eta_{1\xi}\eta_{1\xi\xi}, \quad (\text{c})$$

where

$$\bar{c}_2^2 = T_0/\rho_w,$$

$$\delta = \frac{C_D\rho}{T_0}dV^2 \cos \alpha \sin \alpha,$$

$$\gamma = \frac{C_D\rho}{T_0}dV^2 \sin \alpha,$$

as the differential equations governing the motion of the submerged cable. The constant c_1 is the velocity of propagation of a longitudinal wave in the cable, while the constants c_2 and $\bar{c}_2$ represent the propagation velocities of a transverse wave in air and water respectively.

D.3 *Solution of the Perturbation Equations*

We write

$$p(0, t) = P_0(t) + P_1(t),$$

$$q(0, t) = Q_0(t),$$

and take as boundary conditions

$$\begin{aligned}
 p_1(0, t) &= \frac{P_0(t) + P_1(t)}{e}, \quad (a) \\
 p_2(0, t) &= 0, \quad (b) \\
 q_1(0, t) &= Q_0/e. \quad (c)
 \end{aligned}
 \tag{88}$$

That is, we apportion all of the longitudinal boundary motion to p_1 , and all of the transverse boundary motion to q_1 . Equations [(84b), (86b) and (87b)] then give the complete tension due to the longitudinal component of ship motion to first order. As mentioned in the text, this tension is easily obtained from standard references, and is also the greater part of the ship motion tension.

To determine the tensions due to transverse ship motion, we solve (86c) and (87c) for boundary conditions (88b) and (88c). In addition, we have the transition conditions

$$\begin{aligned}
 q_1(L, t) &= \eta_1(0, t), \quad (a) \\
 q_{1x}(L, t) &= \eta_{1\zeta}(0, t), \quad (b) \\
 p_2(x = L, t) &= p_2(\zeta = 0, t), \quad (c) \\
 p_{2x}(L, t) &= p_{2\zeta}(0, t), \quad (d)
 \end{aligned}
 \tag{89}$$

which follow if we assume that at the point of entry into the water the cable is continuous and the tensions are finite and continuous.

We consider only the problem of the tensions associated with a harmonic steady-state transverse disturbance. Equations (86a) and (87a) show the transverse response to this disturbance to be independent of the longitudinal motion to first order. The first-order transverse motion in turn can be thought of as a forcing action on the second order longitudinal motion, as (86c) and (87c) indicate. This suggests the program we follow to compute tensions. Namely, we first determine the first-order steady-state transverse response, then the second-order steady-state longitudinal response which is excited by the first-order transverse oscillation, and finally, by (84c) the resulting tension caused by transverse motion.

D.4 *Transverse Response*

At the ship we assume a harmonic forcing function

$$Q_0(t) = A \cos \omega t, \tag{90}$$

and we introduce complex exponential representation

$$q_1 = \operatorname{Re} Q_1(x) e^{i\omega t},$$

$$\eta_1 = \operatorname{Re} H_1(\zeta) e^{i\omega t},$$

where the factor $e^{i\omega t}$ will be henceforth suppressed.

The solution of (86a) and (87a) for the steady state may then be written

$$Q_1(x) = B_1 \exp \frac{i\omega x}{c_2} + B_2 \exp \left(- \frac{i\omega x}{c_2} \right),$$

$$H_1(\zeta) = F_1 \exp (q_1 \zeta) + F_2 \exp (q_2 \zeta),$$

where the B 's and F 's are complex constants and q_1 and q_2 are the roots of the quadratic

$$q^2 - \delta q - i\omega\gamma + \frac{\omega^2}{\bar{c}_2^2} = 0.$$

Throwing away the root of this equation which corresponds to the incoming wave in water, we get

$$H_1(\zeta) = F \exp (q_1 \zeta).$$

where q_1 is the root corresponding to the outgoing wave. The three complex constants B_1 , B_2 , and F can now be determined from (89a), (89b) and (90)

$$B_1 + B_2 = A/e,$$

$$B_1 \exp \frac{i\omega L}{c_2} + B_2 \exp \left(- \frac{i\omega L}{c_2} \right) - F = 0, \quad (91)$$

$$\frac{i\omega}{c_2} \left[B_1 \exp \frac{i\omega L}{c_2} - B_2 \exp \left(- \frac{i\omega L}{c_2} \right) \right] - q_1 F = 0.$$

We note that B_1 , B_2 and F are proportional to the amplitude A of the forcing motion.

D.5 Second-Order Longitudinal Response

From the preceding results, the right-hand sides of the equations of longitudinal motion (86c) and (87c) can be computed. This computation for (86c) results in

$$\begin{aligned}
\frac{1}{c_1^2} v_{1tt} v_{1x} - v_{1xx} v_{1x} &= \frac{1}{2} \left(\frac{c_2^2}{c_1^2} - 1 \right) \left(\frac{\omega}{c_2} \right)^3 \\
&\times \left[\left(r_1 \sin \frac{2\omega x}{c_2} + r_2 \cos \frac{2\omega x}{c_2} \right) \cos 2\omega t \right. \\
&+ \left(r_3 \cos \frac{2\omega x}{c_2} - r_4 \sin \frac{2\omega x}{c_2} \right) \sin 2\omega t \\
&\left. + r_5 \sin \frac{2\omega x}{c_2} + r_6 \cos \frac{2\omega x}{c_2} \right],
\end{aligned} \tag{92}$$

where the r 's, which are proportional to the square of the amplitude A , are

$$r_1 = \operatorname{Re} (B_1^2 + B_2^2),$$

$$r_2 = \operatorname{Im} (B_1^2 - B_2^2),$$

$$r_3 = \operatorname{Re} (B_1^2 - B_2^2),$$

$$r_4 = \operatorname{Im} (B_1^2 + B_2^2),$$

$$r_5 = 2\operatorname{Re} B_1 \bar{B}_2,$$

$$r_6 = r_5 + \frac{e^2 r_2^2}{A^2}.$$

[The quantity r_6 will be used subsequently.] Similarly, for the right-hand side of (87c) we get

$$\begin{aligned}
\eta_{1\zeta} \left(\frac{1}{c_1^2} \eta_{1tt} - \eta_{1\zeta\zeta} \right) &= e^{-2\rho\zeta} [(a_1 \cos 2\sigma\zeta + a_2 \sin 2\sigma\zeta) \cos 2\omega t \\
&+ (a_1 \sin 2\sigma\zeta - a_2 \cos 2\sigma\zeta) \sin 2\omega t + a_3],
\end{aligned} \tag{93}$$

where

$$q_1 = -(\rho + i\sigma),$$

$$a_1 = \frac{|F|^2}{2} |q_1| \left[\left(\frac{\omega}{c_1} \right)^2 \cos (2f + g) + |q_1|^2 \cos (2f + 3g) \right],$$

$$a_2 = \frac{|F|^2}{2} |q_1| \left[\left(\frac{\omega}{c_1} \right)^2 \sin (2f + g) + |q_1|^2 \sin (2f + 3g) \right],$$

$$a_3 = |F| \left[\left(\frac{\omega}{c_1} \right)^2 + |q_1|^2 \right] \rho,$$

and

$$f = \arg F, \quad 0 \leq f \leq 2\pi,$$

$$g = \arg (-q_1), \quad 0 \leq g \leq \frac{\pi}{2}.$$

It is seen that expression (92) and (93) have terms of the form

$$F(x) \begin{cases} \sin 2\omega t \\ \cos 2\omega t \end{cases} \quad (94)$$

in addition to functions of x (or ζ) alone. In accordance with the idea that the first order transverse motion is a forcing action on the second order longitudinal motion, we take as solutions of (86c) and (87c) functions of the form

$$G(x) \begin{cases} \sin 2\omega t \\ \cos 2\omega t \end{cases}$$

to correspond to terms of the type given by (94) and functions of x (or ζ) alone to correspond to forcing terms which are independent of time. This again gives linear differential equations which can be readily solved. For example, corresponding to the first term in (92) multiplying $\cos 2\omega t$ we have the assumed solution

$$G(x) \cos 2\omega t,$$

and the differential equation

$$\frac{d^2 G}{dx^2} + \frac{4\omega^2}{c_1^2} G = \frac{1}{2} \left[\left(\frac{c_2}{c_1} \right)^2 - 1 \right] \left(\frac{\omega}{c_2} \right)^3 \left[r_1 \sin \frac{2\omega x}{c_2} + r_2 \cos \frac{2\omega x}{c_2} \right].$$

This has the solution

$$G = \left[A_1 \cos \frac{\omega x}{c_1} + A_2 \sin \frac{\omega x}{c_1} + \frac{\omega}{8c_2} \left(r_1 \sin \frac{2\omega x}{c_2} + r_2 \cos \frac{2\omega x}{c_2} \right) \right],$$

where A_1 and A_2 are undetermined constants.

In this manner, the solution for the longitudinal motion can be obtained in terms of a set of constants. These in turn can be evaluated by means of the boundary and transition conditions on p_2 . This evaluation, although straightforward, is very tedious. We shall omit the details of it here. The final result is

Air:

$$e^2 T_2 = \frac{e^2 \omega^2 \overline{EA}}{4c_1 c_2} \left\{ \left[r_2 \sin \frac{2\omega x}{c_1} + \left(r_3 + \frac{c_1}{c_2} r_4 \sin \frac{2\omega L}{c_1} - \frac{c_1}{c_2} r_6 \cos \frac{2\omega L}{c_1} \right) \cos \frac{2\omega x}{c_1} + \frac{c_1}{c_2} r_6 \right] \cos 2\omega t + \left[r_3 \sin \frac{2\omega L}{c_1} - \left(r_2 + \frac{c_1}{c_2} r_4 \cos \frac{2\omega L}{c_1} + \frac{c_1}{c_2} r_6 \sin \frac{2\omega L}{c_1} \right) \cos \frac{2\omega x}{c_1} + \frac{c_1}{c_2} r_4 \right] \sin 2\omega t + \frac{c_2}{c_1} \left[r_5 \left(\cos \frac{2\omega x}{c_2} - \cos \frac{2\omega L}{c_2} \right) - r_2 \left(\sin \frac{2\omega x}{c_2} - \sin \frac{2\omega L}{c_2} \right) - \frac{|F|^2}{4} \right] \right\},$$

Water:

$$e^2 T_2 = \frac{e^2 \omega^2 \overline{EA}}{4c_1 c_2} \left\{ \left[-r_2 \cos \frac{2\omega L}{c_1} + r_3 \sin \frac{2\omega L}{c_1} + \frac{c_1}{c_2} \sin \frac{2\omega L}{c_1} \left(r_4 \sin \frac{2\omega L}{c_1} - r_6 \cos \frac{2\omega L}{c_1} \right) \right] \sin 2\omega \left(t - \frac{\xi}{c_1} \right) + \left[r_2 \sin \frac{2\omega L}{c_1} + r_3 \cos \frac{2\omega L}{c_1} + \frac{c_1}{c_2} \sin \frac{2\omega L}{c_1} \left(r_6 \sin \frac{2\omega L}{c_1} + r_4 \cos \frac{2\omega L}{c_1} \right) \right] \cos 2\omega \left(t - \frac{\xi}{c_1} \right) - |F|^2 \frac{c_2}{c_1} e^{-2\rho\xi} \right\}.$$

D.6 Numerical Results

Since the r 's are each proportional to the square of the amplitude A , the above results indicate that the transverse motion tension varies as A squared also. It is additionally a function of the frequency of ship motion ω , the forward mean ship velocity V , and the stationary tension T_0 . The computation of the transverse ship motion tension for the laying situation was carried out for cable No. 2. The results are shown in Fig. 34. Here we have denoted the transverse motion tension by T_q and have plotted T_q/A^2 against the period of ship motion τ . Rather than the stationary tension T_0 , we have used the depth h , which during laying is directly related to T_0 by $h = T_0/w$. Fig. 34(a) is a plot of T_q/A^2 versus the period $\tau = 2\pi/\omega$ for $h = \frac{1}{2}, 2$ and 3 nautical miles and for $V = 6$ knots. Figure 34(b) is a plot of T_q/A^2 versus τ for $V = 3, 6$, and 9 knots and $h =$ one nautical mile.

For representative laying, for example at 6 knots with a ship period of 6 seconds into a depth of one nautical mile, Fig. 34 gives

$$T_q/A^2 = 0.50 \text{ lb/ft}^2 \text{ (twist unrestrained),}$$

$$T_q/A^2 = 0.93 \text{ lb/ft}^2 \text{ (twist restrained).}$$

For an extreme value of $A = 20$ feet, we get therefore that T_q is between 200 and 370 pounds.

Additionally, by means of the above analysis, one can compute the rate of damping of a transverse disturbance after it enters the water. The results of this computation are shown in Fig. 18 and are discussed in Section 4.1.

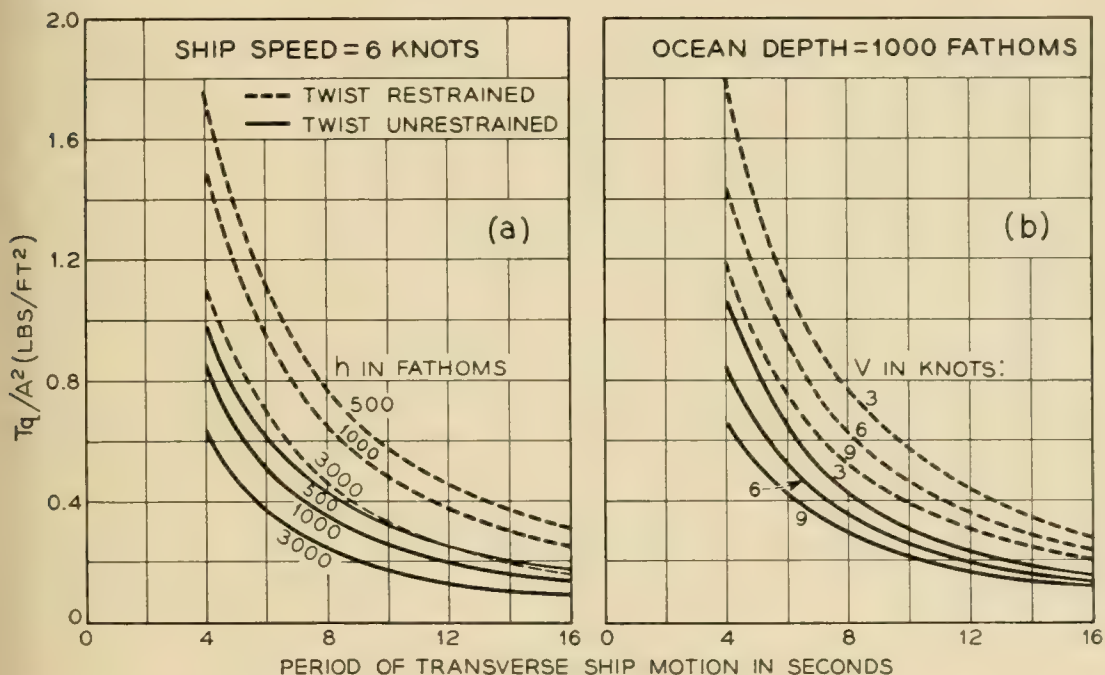


Fig. 34 — Variation of the transverse ship motion tension of cable No. 2 with the period of ship motion.

APPENDIX E

Tension Rise with Time for Suspended Cable

E.1 *Formulation of the Solution of the Problem*

Let O be the lowest point of the cable at time t after the suspension has begun (Fig. 35). We make the following definitions:

h = depth at onset of the suspension,

S_1 = cable length from A to O ,

S_2 = cable length from ship at B to O ,

X_1 = horizontal distance from A to O ,

X_2 = horizontal distance from B to O ,

δ = vertical distance from A to O ,

T_0 = cable tension at O .

If the cable is being paid out with a slack ϵ , then conservation of the total cable length gives the equation

$$S_1 + S_2 = \frac{h}{\sin \alpha} + (1 + \epsilon)Vt + \text{cable stretching.} \quad (95)$$

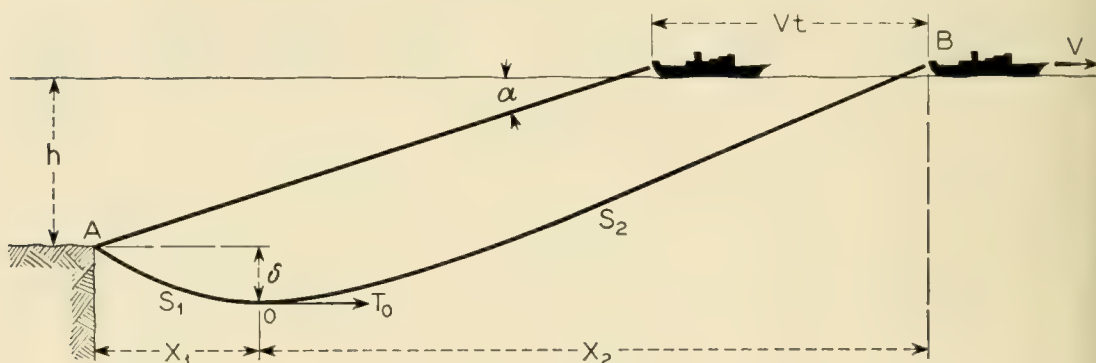


Fig. 35 — Coordinates for the analysis of tension rise when a cable is completely suspended.

It is assumed that there is no cable pulled from the bottom. The cable stretching we evaluate as in the example of Section 3.6, viz.,

$$\text{cable stretching} = (S_1 + S_2) \frac{T_0}{EA}.$$

This makes (95) read

$$S_1 + S_2 = \frac{h}{\sin \alpha} + (1 + \epsilon)Vt + (S_1 + S_2) \frac{T_0}{EA}. \quad (96)$$

To obtain further relations for the unknowns appearing in (96), we assume that from the ship to point O the cable configuration is a stationary one governed by the equations developed in Section 3.6, while from points O to A we assume that the cable configuration is a static catenary. These assumptions yield the following relations:

$$S_1 = \frac{T_0}{w} \sinh \sigma, \quad (a)$$

$$S_2 = \frac{h + \delta}{\sin \alpha} + \kappa \frac{T_0}{w}, \quad (b)$$

$$X_2 = \frac{h + \delta}{\tan \alpha} + \lambda \frac{T_0}{w}, \quad (c)$$

$$\delta = \frac{T_0}{w} (\cosh \sigma - 1), \quad (d) \quad (97)$$

$$X_2 + X_1 = \frac{h}{\tan \alpha} + Vt, \quad (e)$$

$$\sigma = \frac{wX_1}{T_0}, \quad (f)$$

$$T_s = T_0 + w(h + \delta). \quad (g)$$

Here κ and λ are constants, defined and plotted in Section 3.6, which depend only on α , the critical angle corresponding to V . Equations (96) and (97) form a complete set of equations in the unknowns X_1 , X_2 , S_1 , S_2 , T_0 , T_s , δ , and σ . They can be reduced to a set which contain only the unknowns σ and T_s :

$$\varphi_1(\sigma) \sin \alpha - \epsilon \varphi_2(\sigma) \sin \alpha - \bar{h} \left(1 + \frac{\varphi_3(\sigma)}{\varphi_2(\sigma)} \bar{t} \sin \alpha \right) = 0, \quad (a) \quad (98)$$

$$\bar{T}_s = 1 + \frac{\bar{t} \cosh \sigma}{\varphi_2(\sigma)}, \quad (b)$$

where

$$\bar{h} = wh/EA,$$

$$\bar{t} = Vt/h,$$

$$\bar{T}_s = T_s/wh,$$

and

$$\varphi_1 = \sinh \sigma - \sigma + (\cosh \sigma - 1) \tan \frac{\alpha}{2} - (\lambda - \kappa),$$

$$\varphi_2 = \frac{\cosh \sigma - 1}{\tan \alpha} + \lambda + \sigma,$$

$$\varphi_3 = \frac{\cosh \sigma - 1}{\sin \alpha} + \sinh \sigma.$$

A graphical iteration method will be used to solve (98). First we solve

$$\varphi_1(\sigma) \sin \alpha - \epsilon \varphi_2(\sigma) \sin \alpha - \bar{h} = 0 \quad (99)$$

for σ by means of a nomograph to be described later. Next the quantities

$$\frac{\cosh \sigma}{\varphi_2(\sigma)}, \quad \frac{\varphi_3(\sigma)}{\varphi_2(\sigma)} \sin \alpha,$$

are plotted as functions of σ for various α . These plots can then be used as follows to solve (98) for a given $\bar{t}$. Solve (99) to obtain σ_0 . From the plot of $\varphi_3(\sigma)/\varphi_2(\sigma) \sin \alpha$ compute

$$\bar{h}_1^* = \bar{h} \left(1 + \frac{\varphi_3(\sigma_0)}{\varphi_2(\sigma_0)} \bar{t} \sin \alpha \right).$$

Using the value $\bar{h}_1^*$ for $\bar{h}$, compute σ_1 from (99). With this value of σ_1 , compute $\bar{h}_2^*$ from

$$\bar{h}_2^* = \bar{h} \left(1 + \frac{\varphi_3(\sigma_1)}{\varphi_2(\sigma_1)} \bar{t} \sin \alpha \right),$$

etc. In this way a convergent sequence $\sigma_0, \sigma_1, \dots, \sigma_n$ is generated. Finally, from the plot of $\cosh \sigma/\varphi_2(\sigma)$ obtain $\bar{T}_s$.

The above iteration procedure sounds tedious. Actually, in most cases the iteration is not necessary because σ remains essentially independent of time. Thus, the solution of (99) by means of the accompanying nomograph will usually give the complete solution of the problem.

E.2 Nomograph (Alignment Chart) for the Solution of Equation (99)

The relations

$$\begin{aligned} x_1 &= 0, & x_2 &= d, & x_3 &= \frac{10\varphi_2(\sigma) d \sin \alpha}{1 + 10\varphi_2(\sigma) \sin \alpha}, \\ y_1 &= \bar{h}, & y_2 &= \frac{\epsilon}{10}, & y_3 &= \frac{\varphi_1(\sigma) \sin \alpha}{1 + 10\varphi_2(\sigma) \sin \alpha}, \end{aligned} \quad (100)$$

where d is an arbitrary constant define parametrically three curves

$$y_i = y_i(x_i), \quad i = 1, 2, 3,$$

which we imagine plotted on a cartesian (x, y) coordinate system. A set of values $\bar{h}$, ϵ , and σ determine three points (x_i, y_i) ($i = 1, 2, 3$) which lie on these curves. If these points lie on a straight line, it can be shown that they satisfy (99).

On the left-hand sides of Fig. 36 we have plotted the curves given by (100) for various values of the critical angle α . The values of the parameters $\bar{h} = wh/\bar{E}A$ and ϵ , which describe the curves $y_1 = y_1(x_1)$ and $y_2 = y_2(x_2)$ respectively, are plotted on the indicated scales. Rather than indicate the values of σ along the curve $y_3 = y_3(x_3)$ we have for con-

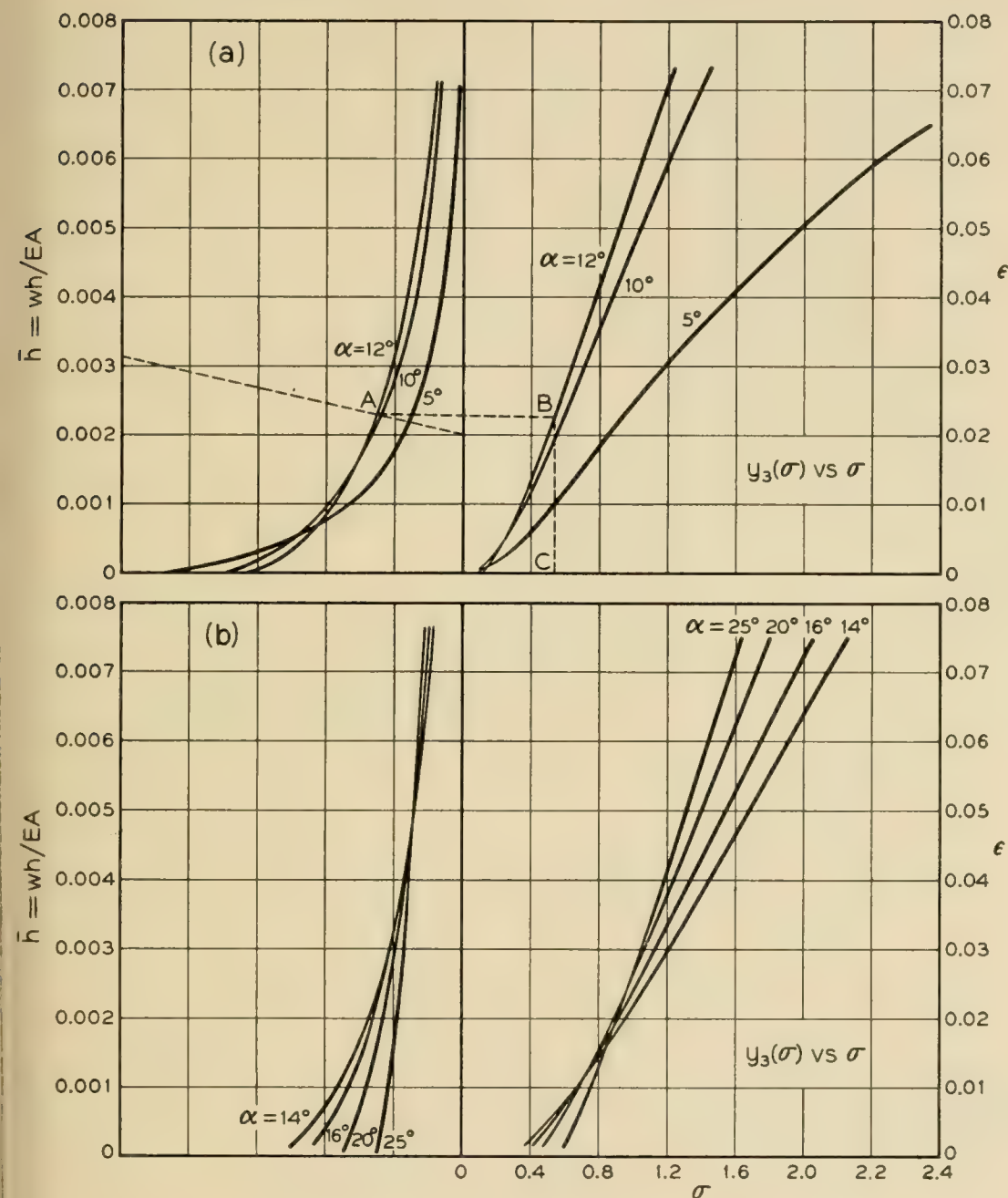


Fig. 36 — Nomograph for the solution of equation (99).

venience made an auxiliary plot on the right-hand sides of Fig. 36 of $y_3(\sigma)$ versus σ , but with numerical values of the ordinate omitted.

In addition, we have plotted in Figs. 37 and 38 the functions $1/\varphi_2 (\cosh \sigma)$ and $(\varphi_3/\varphi_2) \sin \alpha$ for various α .

E.3 Numerical Example

To illustrate the method of obtaining the tension rise with time described above, we consider a numerical example for cable No. 2. The

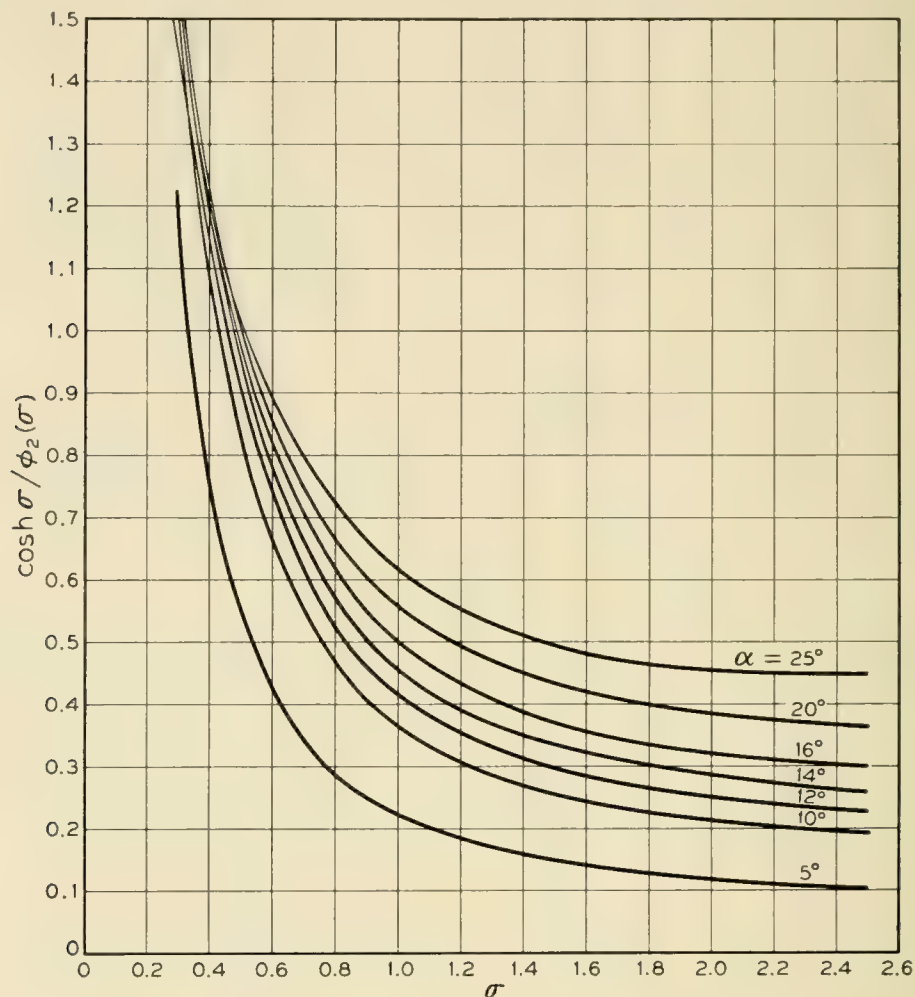


Fig. 37 — Variation of $\frac{\cosh \sigma}{\phi_2(\sigma)}$ with σ .

values we assume for the parameters which enter the calculation are the following:

$$\epsilon = 0.02,$$

$$V = 6 \text{ knots, } (\alpha \approx 12^\circ),$$

$$h = 6,000 \text{ ft,}$$

$$\overline{EA} = 1.2 \times 10^6 \text{ lbs,}$$

$$\bar{h} = 3.1 \times 10^{-3}.$$

To solve (99), we connect on Fig. 36 the points $\epsilon = 0.02$ and $\bar{h} = wh/\overline{EA} = 0.0031$ with a straight edge and note the intersection with the intermediate $y_3 = y_3(x_3)$ curve for $\alpha = 12^\circ$ (point A). We then locate the point on the y_3 versus σ curve having the same ordinate (point B).

Finally, we obtain the root of (99), $\sigma = 0.555$ by reading off the corresponding abscissa (point *C*). This value of σ now serves as the starting point in the iteration procedure by which we find the tension corresponding to a given t .

For example, for $\bar{t} = 1.0$ ($t = 600$ seconds) we have the following sequence of values

σ	$(\varphi_3 \sin \alpha)/\varphi_1$	$\bar{h}^* = \bar{h} [1 + (\varphi_3 \sin \alpha) \bar{t}/\varphi_2]$
$\sigma_0 = 0.555$	0.212	0.00379
$\sigma_1 = 0.580$	0.213	0.00380
$\sigma_2 = 0.580$		

with σ converging to the value 0.580. For $\sigma = 0.580$, Fig. 37 gives

$$\cosh \sigma/\varphi_2 = 0.800.$$

Hence, by (98b) for $\bar{t} = 1.0$

$$\bar{T}_s = 1.80,$$

and $T_s = 1.80$ *wh* or 7,600 pounds.

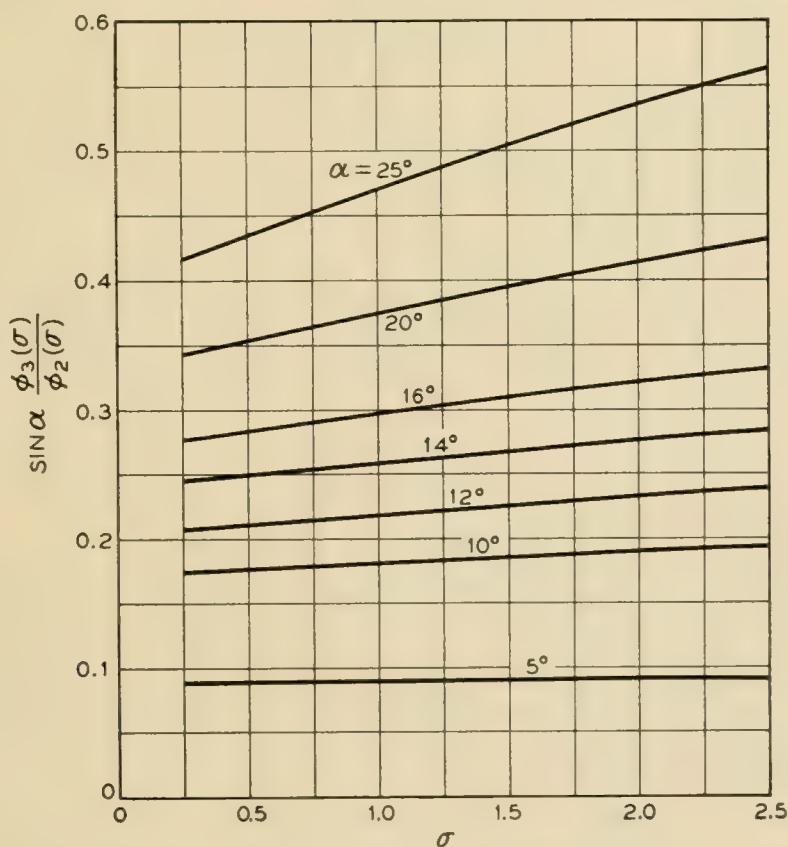


Fig. 38 — Variation of $\sin \alpha \frac{\phi_3(\sigma)}{\phi_2(\sigma)}$ with σ .

APPENDIX F

*The Three-Dimensional Stationary Model*F.1 *Derivation of the Differential Equations*

Let $\vec{i}, \vec{j}, \vec{k}$ be unit vectors along the ξ, η, ζ axes (Fig. 26) and $\vec{t}$ a unit vector along the tangent to the cable configuration in the direction of positive s . As in the two-dimensional model, we take this to be opposite to the direction of travel of the cable elements along the configuration. With respect to the cable configuration the resultant velocity vector of the water is in the $-\vec{i}$ direction. We resolve this velocity into directions normal and tangential to the cable in the plane formed by $\vec{i}$ and $\vec{t}$. The unit vector in the normal direction we denote by $\vec{n}$, namely,

$$\vec{n} = \frac{-\vec{i} + (\vec{i} \cdot \vec{t})\vec{t}}{|-\vec{i} + (\vec{i} \cdot \vec{t})\vec{t}|}. \quad (101)$$

In analogy to the two-dimensional model we assume the normal and tangential drag forces depend only on the corresponding water velocity components. Thus, we take

$$D_N = \frac{C_D \rho d}{2} (\vec{i} \cdot \vec{n} V)^2. \quad (102)$$

Equilibrium of the forces acting on a cable element yields the equation

$$T \frac{d\vec{t}}{ds} + \vec{t} \frac{dT}{ds} + \vec{t} D_T + \vec{n} D_N - \vec{j} w = \rho_c \vec{a}. \quad (103)$$

The vector $\vec{a}$ denotes the acceleration of an element of the cable as it moves at the constant pay-out velocity V_c along the cable configuration. It is easily shown that

$$\vec{a} = V_c^2 \frac{d\vec{t}}{ds}. \quad (104)$$

For convenience we introduce a second reference triad of orthogonal unit vectors $\vec{t}, \vec{u},$ and $\vec{v}$ as follows. The $\vec{v}$ vector is taken in the (ξ, ζ) plane normal to $\vec{t}$; the $\vec{u}$ vector is chosen equal to the vector product $\vec{v} \times \vec{t}$. The angles ψ and θ shown in Fig. 27 describe the orientation of the $(\vec{t}, \vec{u}, \vec{v})$ triad. In terms of these angles, we read from Fig. 27 the following table of direction cosines

	$\vec{i}$	$\vec{j}$	$\vec{k}$
$\vec{t}$	$\cos \theta \cos \psi$	$\sin \theta$	$-\cos \theta \sin \psi$
$\vec{u}$	$-\sin \theta \cos \psi$	$\cos \theta$	$\sin \theta \sin \psi$
$\vec{v}$	$\sin \psi$	0	$\cos \psi$

In the $(\vec{t}, \vec{u}, \vec{v})$ system the vector $\vec{n}$ becomes for example

$$\vec{n} = \frac{\vec{u} \sin \theta \cos \psi - \vec{v} \sin \psi}{[\sin^2 \theta \cos^2 \psi + \sin^2 \psi]^{\frac{1}{2}}}.$$

Imagine the origin of the $(\vec{t}, \vec{u}, \vec{v})$ triad to traverse the cable at unit velocity. The triad during this traverse rotates like a rigid body with respect to the fixed (ξ, η, ζ) frame. The rotation, which we denote by $\vec{\Omega}$, is seen from Fig. 27 to be

$$\vec{\Omega} = \vec{j} \dot{\psi} + \vec{v} \dot{\theta} = \vec{u} \cos \theta \dot{\psi} + \vec{v} \dot{\theta} + \vec{t} \sin \theta \dot{\psi}.$$

Here the dot denotes differentiation with respect to time, or since $ds/dt = 1$ it may be interpreted as differentiation with respect to distance along the cable. The vector $\vec{t}$ is a fixed vector of constant magnitude in the rotating $(\vec{t}, \vec{u}, \vec{v})$ triad, hence

$$\frac{d\vec{t}}{ds} = \vec{\Omega} \times \vec{t} = \vec{u} \dot{\theta} - \vec{v} \cos \theta \dot{\psi}. \quad (105)$$

From (101), (102), (104), and (105) we obtain for (103)

$$\begin{aligned} (T - \rho_c V_c^2)(\vec{u} \dot{\theta} - \vec{v} \cos \theta \dot{\psi}) + \vec{t} \left(\frac{dT}{ds} + D_T \right) \\ + \frac{C_D \rho d V^2}{2} (\sin^2 \theta \cos^2 \psi + \sin^2 \psi)^{\frac{1}{2}} (\vec{u} \sin \theta \cos \psi - \vec{v} \sin \psi) \quad (106) \\ - w(\vec{u} \cos \theta + \vec{t} \sin \theta) = 0, \end{aligned}$$

which gives the three scalar equations, (47).

Further, let $\vec{r}(s)$ be the cable configuration, i.e.,

$$\vec{r}(s) = \vec{i} \xi(s) + \vec{j} \eta(s) + \vec{k} \zeta(s),$$

where $\xi(s)$, $\eta(s)$, and $\zeta(s)$ are the ξ , η , ζ coordinates of a point s of the cable. Then

$$\vec{t} = \vec{i} \frac{d\xi(s)}{ds} + \vec{j} \frac{d\eta(s)}{ds} + \vec{k} \frac{d\zeta(s)}{ds}.$$

Forming the scalar product of $\vec{t}$ with $\vec{i}$, $\vec{j}$, $\vec{k}$ respectively, we get (48) of Section 7.1.

In a θ, ψ, T space the solution trajectories of (47) are given by the solutions of

$$\begin{aligned} (T - \rho_c V_c^2) \frac{d\theta}{dT} &= \frac{\Lambda (\cos^2 \psi \sin^2 \theta + \sin^2 \psi)^{\frac{1}{2}} \cos \psi \sin \theta - \cos \theta}{D_T/w - \sin \theta}, \\ (T - \rho_c V_c^2) \frac{d\psi}{dT} &= \frac{\Lambda (\cos^2 \psi \sin^2 \theta + \sin^2 \psi)^{\frac{1}{2}} \sin \psi}{\cos \theta (D_T/w - \sin \theta)}. \end{aligned}$$

We see that the trajectories are periodic in both θ and ψ with a period of 2π , and only a single region, say

$$0 \leq \theta \leq 2\pi,$$

$$0 \leq \psi \leq 2\pi,$$

need be considered. It is apparent that the straight lines

$$(1) \psi \equiv 0, \quad \theta \equiv \alpha; \quad (2) \psi \equiv 0, \quad \theta \equiv \alpha + \pi;$$

$$(3) \psi \equiv \pi, \quad \theta \equiv 2\pi - \alpha; \quad (4) \psi \equiv \pi, \quad \theta \equiv \pi - \alpha;$$

are solution trajectories which contain all values of T . Along other solution trajectories in this region one easily verifies that

$$T = \frac{\rho_c V_c^2}{w} \times \left(1 - \exp \int_{\theta_0}^{\theta} \frac{\left(\frac{D_T}{w} - \sin \theta \right) d\theta}{\Lambda [\cos^2 \psi(\theta) \sin^2 \theta + \sin^2 \psi(\theta)]^{\frac{1}{2}} \cos \psi(\theta) \sin \theta - \cos \theta} \right),$$

where $\psi = \psi(\theta)$ is obtained from the solution of

$$\frac{d\psi}{d\theta} = \frac{[\cos \theta - \Lambda(\cos^2 \psi \sin^2 \theta + \sin^2 \psi)^{\frac{1}{2}} \cos \psi \sin \theta] \cos \theta}{\Lambda(\cos^2 \psi \sin^2 \theta + \sin^2 \psi)^{\frac{1}{2}} \sin \psi}.$$

From the definitions of ψ and θ , it follows that the lines (3) and (4) are physically identical with lines (1) and (2), and represent straight-line laying and recovery respectively. Likewise, the expression for T shows that any non-straight line trajectory with zero bottom tension is bounded by $\rho_c V_c^2/w$. Hence, as in the case of the two-dimensional model, we conclude that if the tension is somewhere greater than $\rho_c V_c^2/w$ and the bottom tension is zero, the only possible stationary configuration is the straight line lying in the plane of the resultant ship velocity and gravity vectors, and making the critical angle α with the horizontal.

F.2 Perturbation Solution for a Uniform Cross Current

At the outset we assume the tangential drag force to be zero. This gives by (49)

$$T = w(h + \eta), \quad (107)$$

where h is the total ocean depth. Furthermore, we take $\rho_c V_c^2$ to be zero.

If the angle φ (Fig. 26) is small compared to unity, we assume that θ and ψ will vary only slightly from the values they would have if the

upper, cross-current stratum extended all the way to the ocean bottom. That is, we take θ and ψ to be of the form

$$\begin{aligned}\theta &= \alpha' + \bar{\theta}, \\ \psi &= \tilde{\psi},\end{aligned}\tag{108}$$

where α' is the stationary incidence angle corresponding to the velocity V' , and $\bar{\theta}$ and $\tilde{\psi}$ are assumed small compared to unity.

Substituting (48b), (107), and (108) into (47a, b) and retaining only linear terms in θ , ψ and their derivatives, we get the linear first order equations

$$(h + \eta) \frac{d\bar{\theta}}{d\eta} + (2\text{ctn}^2\alpha' + 1)\bar{\theta} = 0, \tag{a}$$

$$(h + \eta) \frac{d\tilde{\psi}}{d\eta} + \text{csc}^2\alpha' \tilde{\psi} = 0. \tag{b}$$

Because in the lower stratum the cable is a straight line parallel to the path of the ship, we have as boundary conditions:

$$\eta = -h', \begin{cases} \bar{\theta} = \alpha - \alpha', \\ \tilde{\psi} = \varphi, \end{cases} \tag{110}$$

where h' is the depth of the upper, cross-current stratum and α is the stationary incidence angle corresponding to the velocity V .

The solution of (109) for the boundary conditions (110) is

$$\begin{aligned}\bar{\theta} &= \left(\frac{h - h'}{h + \eta}\right)^\mu \Delta\alpha, \\ \tilde{\psi} &= \left(\frac{h - h'}{h + \eta}\right)^\nu \varphi,\end{aligned}\tag{111}$$

where

$$\begin{aligned}\mu &= (2\text{ctn}^2\alpha' + 1), \\ \nu &= \text{csc}^2\alpha', \\ \Delta\alpha &= \alpha - \alpha'.\end{aligned}\tag{112}$$

Equation (48) for the space-coordinates ξ , η , and ζ of the cable in turn can be written to terms of first order in the form

$$\begin{aligned}\frac{d\xi}{d\eta} &= \text{ctn}\alpha' - \bar{\theta}\text{csc}^2\alpha', \\ \frac{d\zeta}{d\eta} &= -\tilde{\psi}\text{ctn}\alpha'.\end{aligned}\tag{113}$$

Substituting (111) into (113) and integrating under the condition that at $\eta = 0$, $\xi = 0$ and $\zeta = 0$, we find

$$\begin{aligned}\xi &= \eta \operatorname{ctn} \alpha' + \frac{(h - h') \Delta \alpha}{(\mu - 1)} \operatorname{csc}^2 \alpha' \left[\frac{1}{(h + \eta)^{\mu-1}} - \frac{1}{h^{\mu-1}} \right], \\ \zeta &= \operatorname{ctn} \alpha' \frac{(h - h')}{(\nu - 1)} \varphi \left[\frac{1}{(h + \eta)^{\nu-1}} - \frac{1}{h^{\nu-1}} \right].\end{aligned}\quad (114)$$

These equations describe the space curve formed by the cable in the cross-current stratum.

To determine the distances d and e (Fig. 32), we transform (114) for the cable configuration to coordinates ξ' and ζ' oriented along the ship's path and normal to it respectively by means of

$$\begin{aligned}\xi' &= \xi \cos \varphi - \zeta \sin \varphi, \\ \zeta' &= \xi \sin \varphi + \zeta \cos \varphi.\end{aligned}$$

The result to terms of the first order is

$$\begin{aligned}\xi' &= \eta \operatorname{ctn} \alpha' + \frac{(h - h') \Delta \alpha}{(\mu - 1)} \operatorname{csc}^2 \alpha' \left[\frac{1}{(h + \eta)^{\mu-1}} - \frac{1}{h^{\mu-1}} \right], \\ \zeta' &= \varphi \left(\eta \operatorname{ctn} \alpha' + \frac{\operatorname{ctn} \alpha' (h - h')}{(\nu - 1)} \left[\frac{1}{(h + \eta)^{\nu-1}} - \frac{1}{h^{\nu-1}} \right] \right).\end{aligned}$$

Letting $\eta = -h'$ and denoting the corresponding values of ξ' and ζ' by $-d$ and $-e$ respectively, we obtain (52).

ACKNOWLEDGMENTS

Space does not allow the author to thank individually the many persons who offered comments, corrections, and information during the preparation of this paper. He is especially grateful however to C. H. Elmendorf and H. N. Uptegrove who instigated the study and provided encouragement and support during its preparation, to R. C. Prim and S. P. Morgan for critical comments on the manuscript, to B. C. Heezen for oceanographic information, to A. G. Norem, R. L. Peek, and J. F. Shea for the use of excellent unpublished earlier work on the problem, to D. Ross for assistance on cable hydrodynamics, and to N. C. Youngstrom for invaluable consultations on practical aspects of the submarine cable art and for his collaboration on the work in Appendix B.

REFERENCES

1. W. H. Thomson (Lord Kelvin), On Machinery for Laying Submarine Telegraph Cables, *The Engineer*, **4**, pp. 185-186, Sept. 11, 1857.

2. W. H. Thomson, Machinery for Laying Submarine Cables, *The Engineer*, **4**, p. 280, Oct. 16, 1857.
3. J. A. Longridge and C. E. Brooks, On Submerging Telegraph Cables, *Proceedings of the Institution of Civil Engineers*, pp. 269-314, Feb. 16, 1858.
4. G. B. Airy, On the Mechanical Conditions of the Deposit of a Submarine Cable, *Philosophical Magazine*, **16**, pp. 1-18, July, 1858.
5. W. Gravatt, On the Atlantic Cable, *Philosophical Magazine*, **16**, pp. 34-37, July, 1858.
6. W. S. B. Woolhouse, On the Deposit of Submarine Cables, *Philosophical Magazine*, **19**, pp. 345-364, May, 1860.
7. W. H. Thomson, On the Forces Concerned in the Laying and Lifting of Deep Sea Cables, *Mathematical and Physical Papers*, Vol. 2, Cambridge Univ. Press, Cambridge, 1884, pp. 153-167.
8. E. J. Routh, *Dynamics of a System of Rigid Bodies*, Vol. 2, Macmillan Co., London, 1905, pp. 401-403.
9. W. Siemens, Contributions to the Theory of Submerging and Testing Submarine Telegraphs, *Journal of the Society of Telegraph Engineers*, **5**, pp. 42-66, Feb. 1875. Reprinted in *Scientific and Technical Papers of Werner von Siemens*, John Murray, London, Vol. I, 1892, pp. 237-263. Also see discussions of Siemens paper in the *J. of the Soc. of Telegraph Eng.*, **5**, by Willoughby Smith, pp. 67-76; F. C. Webb, pp. 92-97; W. Siemens, pp. 100-102; W. H. Preece, pp. 102-106; J. A. Longridge, pp. 107-112; F. C. Webb, pp. 113-115; and W. Siemens, pp. 122-126.
10. I. Brunelli, Note sur la théorie mécanique de l'immersion des câbles sous-marins, *Journal Télégraphique*, **36**, pp. 4-6, 25-29, and 49-53, 1912.
11. L. Pode, An Experimental Investigation of the Hydrodynamic Forces on Stranded Cables, Report **713**, David W. Taylor Model Basin, Navy Department, May 1950.
12. L. Pode, Tables for Computing the Equilibrium Configuration of a Flexible Cable in a Uniform Stream, Report **687**, David W. Taylor Model Basin, Navy Department, March 1951.
13. J. W. Johnson (editor), *Proceedings of the First Conference on Ships and Waves at Hoboken, New Jersey, Oct. 1954*. Published by the Council on Wave Research and the Society of Naval Architects and Marine Engineers, 1955.
14. F. Horn, High Sea Test Trip Oscillation and Acceleration Measurements, Translation No. **26**, U. S. Experimental Model Basin, Navy Yard, Washington, D. C., March 1936.
15. F. Eisner, Das Widerstands Problem, *Third International Congress of Applied Mechanics*, Stockholm, 1930.



Theory of Curved Circular Waveguide Containing an Inhomogeneous Dielectric

By SAMUEL P. MORGAN

(Manuscript received February 25, 1957)

Generalized telegraphist's equations are derived, following Schelkunoff, for all modes in a curved circular waveguide containing an inhomogeneous dielectric. Particular attention is paid to the coupling between the TE_{01} mode and other modes in the curved guide. The results are applied to the problem of preventing the mode conversion from TE_{01} to TM_{11} which normally occurs in a curved round waveguide, by partially filling the cross section of the curved guide with a suitably shaped dielectric, such as polystyrene foam. Design equations are given for various compensators, and criteria are set up for keeping the power levels of all spurious modes low in a compensated bend. Dielectric losses, which may be important at millimeter wavelengths, are briefly treated. The potentialities of different compensator designs are illustrated by numerical examples.

INTRODUCTION

It has been recognized for several years that a major problem in the transmission of circular electric waves through multimode round waveguides is the question of negotiating bends. Theoretical studies^{1, 2, 3} have shown that a gentle bend couples the TE_{01} mode to the TE_{11} , TE_{12} , TE_{13} , $\dots$ modes and to the TM_{11} mode. The TM_{11} mode presents the most serious problem, since it has the same phase velocity as TE_{01} in a perfectly conducting straight guide. It follows that power introduced in the TE_{01} mode at the beginning of a gradual bend will be essentially completely transferred to the TM_{11} mode at odd multiples of a certain critical bending angle ϑ_c . The angle ϑ_c is proportional to the ratio of wavelength to guide radius but independent of bending radius; in other words, power transfer cannot be avoided merely by using a sufficiently gentle bend.

S. E. Miller⁴ has discussed a number of methods for transmitting the circular electric wave around bends with small net power loss to TM_{11} . These methods are of two general types.

In the first type the TE_{01} mode, which is not itself a normal mode of the curved guide, is deliberately converted to a combination of TE_{01} and TM_{11} which is a normal mode, or to a particular polarization of TM_{11} which is another normal mode of the curved guide. After traversing the bend the energy is reconverted to TE_{01} . A disadvantage of the normal-mode approach is that the mode conversions necessary at the ends of the bend are frequency sensitive, so that bandwidths appear to be limited to the order of 10 per cent.

A second approach to the bend problem is to break up, by some modification of the guide, the degeneracy which exists between the propagation constants of the TE_{01} and TM_{11} modes in a perfectly conducting straight pipe. The two modes are still coupled by the curvature of the guide, but as Miller has shown, the maximum power transfer will be small if there is sufficient difference between the phase constants or between the attenuation constants of the coupled modes. A difference in phase constants may be provided, for example, by circular corrugations in the waveguide wall, or perhaps most easily by applying a thin layer of dielectric to the inner surface of the guide.⁵ Differential attenuation may be introduced into the TM_{11} mode by a number of methods, in particular by making the guide out of spaced copper rings or a closely-wound wire helix surrounded by a lossy sheath.⁶ Unfortunately, the larger the guide diameter in wavelengths the more difficult it is to get the separation of propagation constants necessary to negotiate a bend of given radius satisfactorily.

Still another solution of the bend problem is to decouple the TE_{01} and TM_{11} modes in a curved guide by partially filling the cross section of the curved guide with dielectric material. The dielectric must be arranged to produce coupling between the TE_{01} and TM_{11} modes which is equal and opposite to the coupling produced by the curvature of the guide. This condition may be satisfied in a great variety of ways; but it is not the only requirement for a good bend compensator. Practical restrictions are that the power levels of all other modes which are coupled to TE_{01} by the dielectric-compensated bend must be kept low, and of course dielectric losses in the compensator must not be excessive.

Part I of this paper treats the general problem of a curved circular waveguide containing an inhomogeneous dielectric. A convenient formulation of the problem is provided by S. A. Schelkunoff's generalized telegraphist's equations for waveguides.⁷ The field at any cross section of the dielectric-compensated curved circular guide is represented as a superposition of the fields of the normal modes of an air-filled straight circular guide. A current amplitude and a voltage amplitude are asso-

ciated with each normal mode, and the currents and voltages satisfy an infinite set of generalized telegraphist's equations. The coupling terms in these equations depend upon the curvature of the guide axis and upon the variation of dielectric permittivity over the cross section. In the present application, the distribution of dielectric is taken to be independent of distance along the bend.*

The generalized telegraphist's equations for all modes in a curved circular waveguide containing an inhomogeneous dielectric are set up in Section 1.1. As a special case one has the equations for an air-filled curved guide, or for a straight guide with an inhomogeneous dielectric. In Section 1.2 we transform from current and voltage amplitudes to the amplitudes of forward and backward traveling waves on a system of coupled transmission lines and in Section 1.3 we work out in some detail the coupling coefficients which involve the TE_{01} mode. An approximate formula for dielectric loss in a compensated bend, which is valid at least in the important practical case when the relative permittivity of the dielectric differs but little from unity, is given in Section 1.4.

Part II applies the foregoing theory to the design of bend compensators for the TE_{01} mode. In a well-designed compensator the amplitudes of the backward (reflected) waves are very low, so we shall neglect reflections. The amplitudes of the spurious forward waves should also be low compared to TE_{01} , so that we may consider them one at a time. We assume that the TE_{01} mode crosstalks independently into each spurious mode, and represent the interaction between modes by a pair of linear, first-order, differential equations in the wave amplitudes. Miller's treatment⁸ of these equations is reviewed in Section 2.1, and applied in Section 2.2 to TE_{01} - TM_{11} coupling in plain and compensated bends. Some results of the Jouguet-Rice theory^{1, 2} for plain bends are confirmed by coupled-line theory. The condition for decoupling TE_{01} and TM_{11} in a compensated bend is written down, and the consequences of imperfect decoupling are discussed.

Three different compensator designs are described in Section 2.3, and evaluated with regard to mode conversions and approximate dielectric losses. In the first case, which may be called the "geometrical optics" solution, the permittivity is supposed to vary continuously in such a way that a bundle of parallel rays entering the bend would be bent into co-axial circular arcs all of the same optical length. This is not a perfect solution of the problem if the wavelength is finite, but it is of some

* We shall not consider the effects of random inhomogeneities, such as bubbles in polystyrene foam, although these might conceivably add to the mode conversion if their dimensions were comparable to the operating wavelength.

theoretical interest nevertheless. In the second case, a dielectric sector of constant permittivity is attached to the inner surface of the bend nearest the center of curvature; the angle of the sector is determined to satisfy the decoupling condition. Such a sector may be an effective compensator if the guide is small enough to propagate only 40 to 50 modes at the operating frequency, as, for example, a $\frac{7}{8}$ -inch guide at a wavelength of 5.4 mm. Finally, we consider a compensator made of three dielectric sectors, whose angles and spacings are chosen to decouple the modes (TE_{21} and TE_{31}) with phase velocities closest to TE_{01} . The three-sector compensator may be necessary if the guide is large enough to propagate 200 to 300 modes, say a 2-inch guide at 5.4 mm.

In Section 2.4 we investigate the effect of increasing the attenuation of the spurious modes generated by the compensated bend. The conclusion is that it is not feasible to add enough loss to the worst spurious modes to reduce appreciably the power which they abstract from TE_{01} .

As a sample of numerical results, it appears possible to negotiate a 90° bend of radius 20 inches in the $\frac{7}{8}$ -inch guide at 5.4 mm with an insertion loss of about 0.3 db. This assumes a single-sector polyfoam compensator with a relative permittivity of 1.036 and a loss tangent of 5×10^{-5} (polyfoam with approximately these constants is currently available). About 0.2 db of the quoted loss is due to mode conversions and 0.1 db to dielectric dissipation. For a 2-inch guide with a three-sector polyfoam compensator, a bending radius of about 12 feet appears feasible. The total loss in a 90° bend should be about 0.35 db, with approximately 0.1 db going into mode conversion and about 0.25 db into dielectric dissipation. The dielectric loss is proportional, of course, to the total bend angle, and for a 180° bend would be double the above figures.

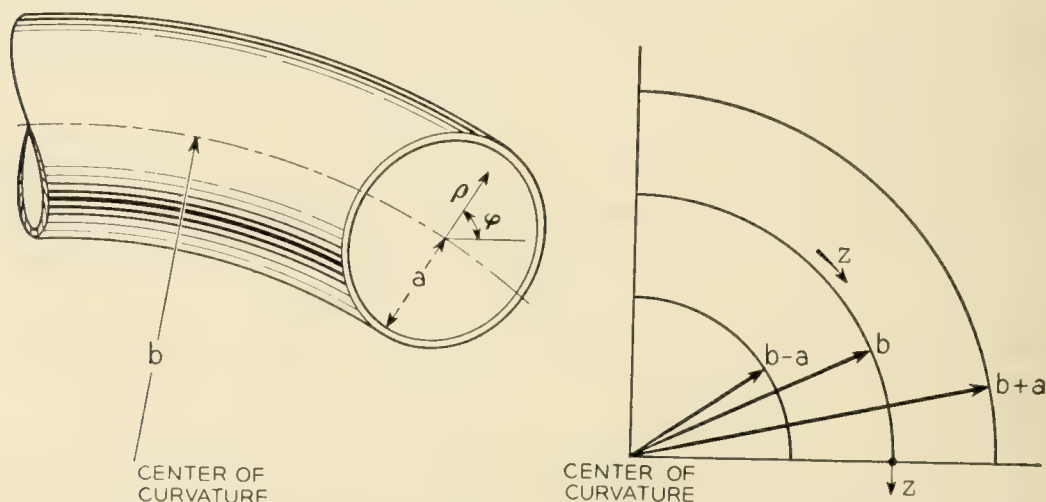


Fig. 1 — Coordinates used in circular bend in circular waveguide.

The appendix contains brief descriptions of three dielectric compensators which can be inserted in a straight section of guide adjacent to a bend. The first two are transducers which convert TE_{01} to a normal mode of the curved guide; they are subject to bandwidth limitations as mentioned by Miller.⁹ The third type merely takes the output mixture of TE_{01} and TM_{11} from a plain bend with a pure TE_{01} input, and reconverts it all to TE_{01} ; it is essentially a broadband device. The spurious modes generated by a bend plus compensator have not been calculated, but it is very unlikely that a smaller bending radius will be permitted when the compensator is outside the bend than when it is inside.

I. THEORY

1.1 Generalized Telegraphist's Equations

To describe electromagnetic fields in a curved circular waveguide one is naturally led to use "bent cylindrical coordinates" (ρ, φ, z) ,^{1, 2, 3} in which the longitudinal coordinate z is distance measured along the curved axis of the guide, while ρ and φ are polar coordinates in a plane normal to the axis of the guide, with origin at the guide axis. The lines $\varphi = 0$ and $\varphi = \pi$ lie in the plane of the bend. The radius of the guide is denoted by a and the radius of the bend (i.e., the radius of curvature of the guide axis) is denoted by b . The coordinate system is shown in Fig. 1.

For the moment regarding (ρ, φ, z) as general orthogonal curvilinear coordinates (u, v, w) , we let

$$u = \rho, \quad v = \varphi, \quad w = z. \quad (1)$$

The element of length in this system is

$$ds^2 = e_1^2 du^2 + e_2^2 dv^2 + e_3^2 dw^2, \quad (2)$$

where

$$e_1 = 1, \quad e_2 = \rho, \quad e_3 = 1 + \xi, \quad (3)$$

and

$$\xi = (\rho/b) \cos \varphi. \quad (4)$$

Maxwell's equations for a field with time dependence $e^{i\omega t}$ may be

written in the form:¹⁰

$$\begin{aligned}
 \frac{1}{e_2 e_3} \left[\frac{\partial}{\partial v} (e_3 E_w) - \frac{\partial}{\partial w} (e_2 E_v) \right] &= -i\omega\mu H_u, \\
 \frac{1}{e_3 e_1} \left[\frac{\partial}{\partial w} (e_1 E_u) - \frac{\partial}{\partial u} (e_3 E_w) \right] &= -i\omega\mu H_v, \\
 \frac{1}{e_1 e_2} \left[\frac{\partial}{\partial u} (e_2 E_v) - \frac{\partial}{\partial v} (e_1 E_u) \right] &= -i\omega\mu H_w, \\
 \frac{1}{e_2 e_3} \left[\frac{\partial}{\partial v} (e_3 H_w) - \frac{\partial}{\partial w} (e_2 H_v) \right] &= i\omega\epsilon E_u, \\
 \frac{1}{e_3 e_1} \left[\frac{\partial}{\partial w} (e_1 H_u) - \frac{\partial}{\partial u} (e_3 H_w) \right] &= i\omega\epsilon E_v, \\
 \frac{1}{e_1 e_2} \left[\frac{\partial}{\partial u} (e_2 H_v) - \frac{\partial}{\partial v} (e_1 H_u) \right] &= i\omega\epsilon E_w.
 \end{aligned} \tag{5}$$

In these equations the permeability μ and the permittivity ϵ may be functions of position. If there is dissipation in the medium, either ϵ or μ or both may be complex.

To convert Maxwell's equations into generalized telegraphist's equations, we introduce the field distributions characteristic of the normal modes of a straight, cylindrical guide filled with a homogeneous dielectric. The derivation follows very closely that given by Schelkunoff¹¹ for an inhomogeneously-filled straight guide. Each mode is described by a transverse field distribution pattern $T(u, v)$, where $T(u, v)$ satisfies

$$\nabla^2 T = \frac{1}{e_1 e_2} \left[\frac{\partial}{\partial u} \left(\frac{e_2}{e_1} \frac{\partial T}{\partial u} \right) + \frac{\partial}{\partial v} \left(\frac{e_1}{e_2} \frac{\partial T}{\partial v} \right) \right] = -\chi^2 T, \tag{6}$$

and χ is a separation constant which takes on discrete values for the various TE and TM modes. We shall denote the function corresponding to the n th TE mode by $T_{[n]}(u, v)$, and the separation constant by $\chi_{[n]}$, with the subscript in *brackets*. The normal derivative of $T_{[n]}$ vanishes on the perfectly conducting waveguide boundary. Similarly the function corresponding to the n th TM mode is denoted by $T_{(n)}(u, v)$, and the separation constant by $\chi_{(n)}$, with the subscript in *parentheses*. The function $T_{(n)}$ vanishes on the boundary of the waveguide. For the present the

modes may be assumed to be numbered in order of increasing cutoff frequency; later when we have occasion to refer to specific modes we shall replace the subscript n by the customary double subscript notation for TE and TM modes in a circular guide.

The T -functions are assumed to be so normalized that

$$\begin{aligned} \int_S (\text{grad } T) \cdot (\text{grad } T) \, dS &\equiv \int_S (\text{flux } T) \cdot (\text{flux } T) \, dS \\ &\equiv \chi^2 \int_S T^2 \, dS = 1, \end{aligned} \tag{7}$$

where S is the cross section of the guide. The gradient and flux of a scalar point-function W are transverse vectors with the following components:

$$\begin{aligned} \text{grad}_u W &= \frac{\partial W}{e_1 \partial u}, & \text{grad}_v W &= \frac{\partial W}{e_2 \partial v}, \\ \text{flux}_u W &= \frac{\partial W}{e_2 \partial v}, & \text{flux}_v W &= -\frac{\partial W}{e_1 \partial u}. \end{aligned} \tag{8}$$

Various orthogonality relationships exist among the T -functions corresponding to different modes of the guide, and among their gradients and fluxes. These relationships have been listed by Schelkunoff.¹²

The transverse components of the fields in the curved guide may be derived from potential and stream functions, U and Ψ for TE waves and V and Π for TM waves. Thus

$$\begin{aligned} E_t &= -\text{grad } V - \text{flux } \Psi, \\ H_t &= \text{flux } \Pi - \text{grad } U. \end{aligned} \tag{9}$$

We now assume series expansions for the potential and stream functions in terms of the functions $T(u, v)$, with coefficients depending on w . Let

$$\begin{aligned} V &= -\sum_n V_{(n)}(w) \, T_{(n)}(u, v), & \Pi &= -\sum_n I_{(n)}(w) \, T_{(n)}(u, v), \\ \Psi &= -\sum_n V_{[n]}(w) \, T_{[n]}(u, v), & U &= -\sum_n I_{[n]}(w) \, T_{[n]}(u, v). \end{aligned} \tag{10}$$

The I 's and V 's have the dimensions and properties of transmission-line currents and voltages. If we substitute (10) into (9) and expand in

components by (8), we obtain:

$$\begin{aligned} E_u &= \sum_n \left[V_{(n)} \frac{\partial T_{(n)}}{e_1 \partial u} + V_{[n]} \frac{\partial T_{[n]}}{e_2 \partial v} \right], \\ E_v &= \sum_n \left[V_{(n)} \frac{\partial T_{(n)}}{e_2 \partial v} - V_{[n]} \frac{\partial T_{[n]}}{e_1 \partial u} \right], \\ H_u &= \sum_n \left[-I_{(n)} \frac{\partial T_{(n)}}{e_2 \partial v} + I_{[n]} \frac{\partial T_{[n]}}{e_1 \partial u} \right], \\ H_v &= \sum_n \left[I_{(n)} \frac{\partial T_{(n)}}{e_1 \partial u} + I_{[n]} \frac{\partial T_{[n]}}{e_2 \partial v} \right]. \end{aligned} \quad (11)$$

For the longitudinal field components it is convenient to expand the combinations $e_3 E_w$ and $e_3 H_w$ in the following series:

$$\begin{aligned} e_3 E_w &= \sum_n \chi_{(n)} V_{w,(n)}(w) T_{(n)}(u, v), \\ e_3 H_w &= \sum_n \chi_{[n]} I_{w,[n]}(w) T_{[n]}(u, v). \end{aligned} \quad (12)$$

It should be noted that the boundary conditions in the curved circular guide are

$$E_v = E_w = 0 \quad (13)$$

at the boundary of the guide, and that these conditions are satisfied by the individual terms of the series for E_v and E_w , on account of the boundary conditions already imposed on $T_{(n)}$ and $T_{[n]}$. Hence we do not have the problem of nonuniform convergence which sometimes arises in treating waveguides of varying cross section by the present method.

The procedure for transforming Maxwell's equations into generalized telegraphist's equations is now straightforward, if a trifle tedious. One substitutes the series (11) and (12) into equations (5), and integrates certain combinations of the latter equations over the cross section of the guide, taking account of the orthogonality properties of the T -functions. For example, subtracting $\partial T_{(m)}/e_2 \partial v$ times the first of equations (5) from $\partial T_{(m)}/e_1 \partial u$ times the second equation, and integrating over the cross section, yields

$$\begin{aligned} \frac{dV_{(m)}}{dw} - \chi_{(m)} V_{w,(m)} &= -i\omega \sum_n \left[I_{(n)} \int_S \mu e_3 (\text{grad } T_{(n)}) \cdot (\text{grad } T_{(m)}) dS \right. \\ &\quad \left. + I_{[n]} \int_S \mu e, (\text{grad } T_{(m)}) \cdot (\text{flux } T_{[n]}) dS \right]. \end{aligned} \quad (14)$$

Adding $\partial T_{[m]}/e_1 \partial u$ times the first equation to $\partial T_{[m]}/e_2 \partial v$ times the second

and integrating gives

$$\begin{aligned} \frac{dV_{[m]}}{dw} = & -i\omega \sum_n \left[I_{(n)} \int_S \mu e_3 (\text{grad } T_{(n)}) \cdot (\text{flux } T_{[m]}) dS \right. \\ & \left. + I_{[n]} \int_S \mu e_3 (\text{grad } T_{[m]}) \cdot (\text{grad } T_{[n]}) dS \right]. \end{aligned} \quad (15)$$

Similarly, from the fourth and fifth equations we get

$$\begin{aligned} \frac{dI_{(m)}}{dw} = & -i\omega \sum_n \left[V_{(n)} \int_S \epsilon e_3 (\text{grad } T_{(m)}) \cdot (\text{grad } T_{(n)}) dS \right. \\ & \left. + V_{[n]} \int_S \epsilon e_3 (\text{grad } T_{(m)}) \cdot (\text{flux } T_{[n]}) dS \right], \end{aligned} \quad (16)$$

$$\begin{aligned} \frac{dI_{[m]}}{dw} - \chi_{[m]} I_{w,[m]} \\ = & -i\omega \sum_n \left[V_{(n)} \int_S \epsilon e_3 (\text{grad } T_{(n)}) \cdot (\text{flux } T_{[m]}) dS \right. \\ & \left. - V_{[n]} \int_S \epsilon e_3 (\text{grad } T_{[m]}) \cdot (\text{grad } T_{[n]}) dS \right]. \end{aligned} \quad (17)$$

From the third of equations (5) using the fact that the T -functions satisfy (6), then multiplying through by $T_{[m]}$ and integrating over the cross section, we get,

$$V_{[m]} = -i\omega \sum_n I_{w,[n]} \chi_{[n]} \int_S \frac{\mu T_{[n]} T_{[m]}}{e_3} dS. \quad (18)$$

Similarly, from the sixth equation,

$$I_{(m)} = -i\omega \sum_n V_{w,(n)} \chi_{(n)} \int_S \frac{\epsilon T_{(n)} T_{(m)}}{e_3} dS. \quad (19)$$

These equations may be written in the form

$$\begin{aligned} V_{[m]} &= -\sum_n \frac{Z_{w,[m][n]} I_{w,[n]}}{\chi_{[m]}}, \\ I_{(m)} &= -\sum_n \frac{Y_{w,(m)(n)} V_{w,(n)}}{\chi_{(m)}}, \end{aligned} \quad (20)$$

where we define

$$\begin{aligned} Z_{w,[m][n]} &= i\omega \chi_{[m]} \chi_{[n]} \int_S \frac{\mu T_{[m]} T_{[n]}}{e_3} dS, \\ Y_{w,(m)(n)} &= i\omega \chi_{(m)} \chi_{(n)} \int_S \frac{\epsilon T_{(m)} T_{(n)}}{e_3} dS. \end{aligned} \quad (21)$$

Solving (20) for $I_{w,[m]}$ and $V_{w,(m)}$ in terms of $V_{[n]}$ and $I_{(n)}$ respectively, we may write

$$\begin{aligned} I_{w,[m]} &= -\sum_n \frac{Y_{w,[m][n]}}{\chi_{[m]}} V_{[n]}, \\ V_{w,(m)} &= -\sum_n \frac{Z_{w,(m)(n)}}{\chi_{(m)}} I_{(n)}, \end{aligned} \quad (22)$$

where $Y_{w,[m][n]}$ and $Z_{w,(m)(n)}$ may be defined as the ratios of certain determinants involving $Z_{w,[m][n]}$ and $Y_{w,(m)(n)}$ respectively.*

We are now able to eliminate $V_{w,(m)}$ and $I_{w,[m]}$ from (14) and (17), and to write down the generalized telegraphist's equations for the curved circular waveguide filled with an inhomogeneous dielectric in the following form:

$$\begin{aligned} \frac{dV_{(m)}}{dw} &= -\sum_n [Z_{(m)(n)} I_{(n)} + Z_{(m)[n]} I_{[n]}], \\ \frac{dV_{[m]}}{dw} &= -\sum_n [Z_{[m](n)} I_{(n)} + Z_{[m][n]} I_{[n]}], \\ \frac{dI_{(m)}}{dw} &= -\sum_n [Y_{(m)(n)} V_{(n)} + Y_{(m)[n]} V_{[n]}], \\ \frac{dI_{[m]}}{dw} &= -\sum_n [Y_{[m](n)} V_{(n)} + Y_{[m][n]} V_{[n]}]. \end{aligned} \quad (23)$$

The impedance and admittance coefficients are defined by:

$$\begin{aligned} Z_{(m)(n)} &= i\omega \int_S \mu e_3 (\text{grad } T_{(m)}) \cdot (\text{grad } T_{(n)}) dS + Z_{w,(m)(n)}, \\ Z_{(m)[n]} &= i\omega \int_S \mu e_3 (\text{grad } T_{(m)}) \cdot (\text{flux } T_{[n]}) dS, \\ Z_{[m](n)} &= i\omega \int_S \mu e_3 (\text{flux } T_{[m]}) \cdot (\text{grad } T_{(n)}) dS, \\ Z_{[m][n]} &= i\omega \int_S \mu e_3 (\text{grad } T_{[m]}) \cdot (\text{grad } T_{[n]}) dS, \\ Y_{(m)(n)} &= i\omega \int_S \epsilon e_3 (\text{grad } T_{(m)}) \cdot (\text{grad } T_{(n)}) dS, \end{aligned} \quad (24)$$

* If (20) are actually regarded as infinite sets of equations, one has to pay some attention to defining the meaning of the statement that their solutions are given by (22). In practice we shall deal with a finite number of modes, and shall bypass all questions about infinite processes.

$$Y_{(m)[n]} = i\omega \int_S \epsilon e_3 (\text{grad } T_{(m)}) \cdot (\text{flux } T_{[n]}) dS,$$

$$Y_{[m](n)} = i\omega \int_S \epsilon e_3 (\text{flux } T_{[m]}) \cdot (\text{grad } T_{(n)}) dS,$$

$$Y_{[m][n]} = i\omega \int_S \epsilon e_3 (\text{grad } T_{[m]}) \cdot (\text{grad } T_{[n]}) dS + Y_{w,[m][n]}.$$

It may be noted that if the guide were straight, so that $e_3 = 1$, and if μ and ϵ were constant over the cross section, the Y 's and Z 's would all be zero except for those having equal subscripts. If the curvature of the guide is gentle, and if μ and ϵ do not vary much over the cross section (or if they vary extensively only in a small part of the cross section), then the Y 's and Z 's with unequal subscripts will be small, and we can obtain approximate solutions of (23) based upon the smallness of the coupling.

We shall now compute first-order approximations to the impedance and admittance coefficients under the following assumptions:

$$\begin{aligned} \mu &= \mu_0, \\ \epsilon &= \epsilon_0(1 + \delta), \end{aligned} \tag{25}$$

where μ_0 and ϵ_0 are constants (usually but not necessarily the permeability and permittivity of free space), and δ is a dimensionless function of position. No mathematical difficulties would follow from assuming μ as well as ϵ to be variable, but since the case of varying permeability is not of such immediate practical interest we shall omit this slight additional complication. In order that the coupling per unit length due to curvature and to inhomogeneity of the dielectric be small, we further assume that

$$\begin{aligned} |\xi| &\leq a/b \ll 1, \\ \frac{1}{S} \int_S |\delta| dS &\ll 1, \end{aligned} \tag{26}$$

where, as usual, S is the cross-sectional area of the guide.

From (21), first-order approximations to $Z_{w,[m][n]}$ and $Y_{w,(m)(n)}$ are

$$\begin{aligned} Z_{w,[m][n]} &= i\omega\mu_0[\delta_{mn} - \xi_{[m][n]}], \\ Y_{w,(m)(n)} &= i\omega\epsilon_0[\delta_{mn} + \delta_{(m)(n)} - \xi_{(m)(n)}], \end{aligned} \tag{27}$$

where

$$\begin{aligned}\xi_{[m][n]} &= \chi_{[m]}\chi_{[n]} \int_S \xi T_{[m]} T_{[n]} dS, \\ \xi_{(m)(n)} &= \chi_{(m)}\chi_{(n)} \int_S \xi T_{(m)} T_{(n)} dS, \\ \delta_{(m)(n)} &= \chi_{(m)}\chi_{(n)} \int_S \delta T_{(m)} T_{(n)} dS.\end{aligned}\tag{28}$$

The quantity δ_{mn} is the Kronecker delta, and is not to be confused with $\delta_{(m)(n)}$, which is defined by the last of equations (28). Note that $\xi_{[m][n]}$ and $\xi_{(m)(n)}$ are zero unless the angular indices of the two modes involved differ by exactly unity.

It is not difficult to obtain approximate solutions of (20) in the forms (22), since the off-diagonal elements of the coefficient matrices of (20) are small compared to the diagonal elements. Using the expression derived by Rice¹³ for the inverse of an almost-diagonal matrix (we shall not attempt to prove this result for infinite matrices), we find the first-order approximations

$$\begin{aligned}Y_{w,[m][n]} &= \frac{\chi_{[m]}\chi_{[n]}}{i\omega\mu_0} [\delta_{mn} + \xi_{[m][n]}], \\ Z_{w,(m)(n)} &= \frac{\chi_{(m)}\chi_{(n)}}{i\omega\epsilon_0} [\delta_{mn} + \xi_{(m)(n)} - \delta_{(m)(n)}].\end{aligned}\tag{29}$$

Approximate expressions for the impedance and admittance coefficients appearing in the generalized telegraphist's equations (23) are:

$$\begin{aligned}Z_{(m)(n)} &= i\omega\mu_0[\delta_{mn} + \Xi_{(m)(n)}] + Z_{w,(m)(n)}, \\ Z_{(m)[n]} &= i\omega\mu_0\Xi_{(m)[n]}, \\ Z_{[m](n)} &= i\omega\mu_0\Xi_{[m](n)}, \\ Z_{[m][n]} &= i\omega\mu_0[\delta_{mn} + \Xi_{[m][n]}], \\ Y_{(m)(n)} &= i\omega\epsilon_0[\delta_{mn} + \Xi_{(m)(n)} + \Delta_{(m)(n)}], \\ Y_{(m)[n]} &= i\omega\epsilon_0[\Xi_{(m)[n]} + \Delta_{(m)[n]}], \\ Y_{[m](n)} &= i\omega\epsilon_0[\Xi_{[m](n)} + \Delta_{[m](n)}], \\ Y_{[m][n]} &= i\omega\epsilon_0[\delta_{mn} + \Xi_{[m][n]} + \Delta_{[m][n]}] + Y_{w,[m][n]}.\end{aligned}\tag{30}$$

where

$$\begin{aligned}
 \Xi_{(m)(n)} &= \int_S \xi (\text{grad } T_{(m)}) \cdot (\text{grad } T_{(n)}) dS, \\
 \Xi_{(m)[n]} &= \int_S \xi (\text{grad } T_{(m)}) \cdot (\text{flux } T_{[n]}) dS, \\
 \Xi_{[m](n)} &= \int_S \xi (\text{flux } T_{[m]}) \cdot (\text{grad } T_{(n)}) dS, \\
 \Xi_{[m][n]} &= \int_S \xi (\text{grad } T_{[m]}) \cdot (\text{grad } T_{[n]}) dS,
 \end{aligned} \tag{31}$$

and

$$\begin{aligned}
 \Delta_{(m)(n)} &= \int_S \delta (\text{grad } T_{(m)}) \cdot (\text{grad } T_{(n)}) dS, \\
 \Delta_{(m)[n]} &= \int_S \delta (\text{grad } T_{(m)}) \cdot (\text{flux } T_{[n]}) dS, \\
 \Delta_{[m](n)} &= \int_S \delta (\text{flux } T_{[m]}) \cdot (\text{grad } T_{(n)}) dS, \\
 \Delta_{[m][n]} &= \int_S \delta (\text{grad } T_{[m]}) \cdot (\text{grad } T_{[n]}) dS.
 \end{aligned} \tag{32}$$

The Ξ 's are zero unless the angular indices of the two modes involved differ by exactly unity.

1.2 Representation in Terms of Coupled Traveling Waves

From now on we shall assume that the distribution of dielectric over the cross section of the curved guide is independent of distance along the guide, so that the impedance and admittance coefficients are constants independent of z . (We shall henceforth designate the coordinates by (ρ, φ, z) , instead of the (u, v, w) of the preceding section.) The generalized telegraphist's equations now represent an infinite set of coupled, uniform transmission lines, and their solution would be equivalent to the solution of an infinite system of linear algebraic equations and the corresponding characteristic equation.

For our purposes it is convenient to write the transmission-line equations not in terms of currents and voltages, but in terms of the amplitudes of forward and backward traveling waves, assumed to exist in a straight guide filled with a homogeneous medium. Thus let a and b be the amplitudes of the forward and backward waves of a typical mode at

a certain point. In what follows the wave amplitudes a and b will always carry mode subscripts, so they need never be confused with guide radius and radius of curvature. The mode current and voltage are related to the wave amplitudes by

$$\begin{aligned} V &= K^{\frac{1}{2}}(a + b), \\ I &= K^{-\frac{1}{2}}(a - b), \end{aligned} \quad (33)$$

where K is the wave impedance. We have

$$\begin{aligned} K_{(n)} &= h_{(n)}/\omega\epsilon_0, \\ K_{[n]} &= \omega\mu_0/h_{[n]}, \end{aligned} \quad (34)$$

for TM and TE waves respectively, where $h_{(n)}$ and $h_{[n]}$ represent the unperturbed phase constants,

$$\begin{aligned} h_{(n)} &= (\beta^2 - \chi_{(n)}^2)^{\frac{1}{2}}, \\ h_{[n]} &= (\beta^2 - \chi_{[n]}^2)^{\frac{1}{2}}, \end{aligned} \quad (35)$$

and

$$\beta^2 = \omega^2\mu_0\epsilon_0. \quad (36)$$

For a cutoff mode, h is negative imaginary; but we shall deal only with propagating modes in the present analysis.

If we represent the currents and voltages in the generalized telegraphist's equations (23) in terms of the traveling-wave amplitudes, and then perform some obvious additions and subtractions, we obtain the following equations for coupled traveling waves:

$$\begin{aligned} \frac{da_{(m)}}{dz} &= -i \sum_n [\kappa_{(m)(n)}^+ a_{(n)} + \kappa_{(m)(n)}^- b_{(n)} + \kappa_{(m)[n]}^+ a_{[n]} + \kappa_{(m)[n]}^- b_{[n]}], \\ \frac{db_{(m)}}{dz} &= +i \sum_n [\kappa_{(m)(n)}^- a_{(n)} + \kappa_{(m)(n)}^+ b_{(n)} + \kappa_{(m)[n]}^- a_{[n]} + \kappa_{(m)[n]}^+ b_{[n]}], \\ \frac{da_{[m]}}{dz} &= -i \sum_n [\kappa_{[m](n)}^+ a_{(n)} + \kappa_{[m](n)}^- b_{(n)} + \kappa_{[m][n]}^+ a_{[n]} + \kappa_{[m][n]}^- b_{[n]}], \\ \frac{db_{[m]}}{dz} &= +i \sum_n [\kappa_{[m](n)}^- a_{(n)} + \kappa_{[m](n)}^+ b_{(n)} + \kappa_{[m][n]}^- a_{[n]} + \kappa_{[m][n]}^+ b_{[n]}]. \end{aligned} \quad (37)$$

The κ 's are coupling coefficients defined in terms of the impedance and admittance coefficients by

$$\begin{aligned} i\kappa_{(m)(n)}^{\pm} &= \frac{1}{2}[(K_{(m)}K_{(n)})^{\frac{1}{2}} Y_{(m)(n)} \pm (K_{(m)}K_{(n)})^{-\frac{1}{2}} Z_{(m)(n)}], \\ i\kappa_{(m)[n]}^{\pm} &= \frac{1}{2}[(K_{(m)}K_{[n]})^{\frac{1}{2}} Y_{(m)[n]} \pm (K_{(m)}K_{[n]})^{-\frac{1}{2}} Z_{(m)[n]}], \\ i\kappa_{[m](n)}^{\pm} &= \frac{1}{2}[(K_{[m]}K_{(n)})^{\frac{1}{2}} Y_{[m](n)} \pm (K_{[m]}K_{(n)})^{-\frac{1}{2}} Z_{[m](n)}], \\ i\kappa_{[m][n]}^{\pm} &= \frac{1}{2}[(K_{[m]}K_{[n]})^{\frac{1}{2}} Y_{[m][n]} \pm (K_{[m]}K_{[n]})^{-\frac{1}{2}} Z_{[m][n]}]. \end{aligned} \quad (38)$$

In these definitions the plus signs are taken together, likewise the minus signs. The factors of i are introduced in order that the κ 's may be real for propagating modes in a lossless guide.

For the small-coupling case discussed at the end of the preceding section, it is convenient to separate a typical coupling coefficient into two parts; thus,

$$\kappa = c + d. \quad (39)$$

Here c is the coupling coefficient due to curvature and is zero unless the angular indices of the two modes involved differ by unity. The coupling coefficient d is due to the dielectric. All d 's vanish if the dielectric is homogeneous; otherwise particular symmetries may cause certain classes of d 's to be zero. The c 's and d 's may be expressed in terms of integrals written down in the preceding section if we substitute for the Y 's and Z 's in (38) their definitions according to (30).

The κ^+ 's which have equal subscripts $(n)(n)$ or $[n][n]$ may be regarded as phase constants (of particular TM or TE modes) which have been modified by the presence of the dielectric. For the modified phase constants we introduce the symbols $\beta_{(n)}$ and $\beta_{[n]}$; thus,

$$\begin{aligned} \beta_{(n)} &\equiv \kappa_{(n)(n)}^+ = h_{(n)} + \frac{\chi_{(n)}^2 \delta_{(n)(n)} + h_{(n)}^2 \Delta_{(n)(n)}}{2h_{(n)}}, \\ \beta_{[n]} &\equiv \kappa_{[n][n]}^+ = h_{[n]} + \frac{\beta^2 \Delta_{[n][n]}}{2h_{[n]}}. \end{aligned} \quad (40)$$

The general expressions for the coupling coefficients between any two different modes are as follows:

$$\begin{aligned} c_{(m)(n)}^\pm &= \frac{1}{2} \left[\sqrt{h_{(m)} h_{(n)}} \Xi_{(m)(n)} \pm \frac{\beta^2 \Xi_{(m)(n)} - \chi_{(m)} \chi_{(n)} \xi_{(m)(n)}}{\sqrt{h_{(m)} h_{(n)}}} \right], \\ d_{(m)(n)}^\pm &= \frac{1}{2} \left[\sqrt{h_{(m)} h_{(n)}} \Delta_{(m)(n)} \pm \frac{\chi_{(m)} \chi_{(n)} \delta_{(m)(n)}}{\sqrt{h_{(m)} h_{(n)}}} \right], \\ c_{(m)[n]}^\pm &= \frac{1}{2} \beta \Xi_{(m)[n]} [\sqrt{h_{(m)}/h_{[n]}} \pm \sqrt{h_{[n]}/h_{(m)}}], \\ d_{(m)[n]}^\pm &= \frac{1}{2} \beta \Delta_{(m)[n]} \sqrt{h_{(m)}/h_{[n]}}, \\ c_{[m](n)}^\pm &= \frac{1}{2} \beta \Xi_{[m](n)} [\sqrt{h_{(n)}/h_{[m]}} \pm \sqrt{h_{[m]}/h_{(n)}}], \\ d_{[m](n)}^\pm &= \frac{1}{2} \beta \Delta_{[m](n)} \sqrt{h_{(n)}/h_{[m]}}, \\ c_{[m][n]}^\pm &= \frac{1}{2} \left[\frac{\beta^2 \Xi_{[m][n]} - \chi_{[m]} \chi_{[n]} \xi_{[m][n]}}{\sqrt{h_{[m]} h_{[n]}}} \pm \Xi_{[m][n]} \sqrt{h_{[m]} h_{[n]}} \right], \\ d_{[m][n]}^\pm &= \frac{\beta^2 \Delta_{[m][n]}}{2 \sqrt{h_{[m]} h_{[n]}}}, \end{aligned} \quad (41)$$

where the symbols with double subscripts are defined by (28), (31), and (32).

1.3 Coupling Coefficients Involving the TE_{01} Mode

In Part II we shall consider a well-compensated bend in which the total power in all spurious modes is everywhere low compared to the power level of TE_{01} . (This is somewhat more restrictive than merely assuming that the power in any one spurious mode is everywhere small compared to the power in TE_{01} .) To first order, therefore, we may compute the power abstracted from TE_{01} by mode conversion by assuming that the TE_{01} mode crosstalks into each spurious mode independently. For this calculation we need the values of the forward coupling coefficients between TE_{01} and all other modes. The crosstalk into backward modes will be negligible in all practical cases.

We shall use the customary double-subscript notation for the various modes in a round guide, but shall continue to denote TM waves with parentheses and TE waves with brackets. *We assume that the distribution of dielectric is symmetric with respect to the plane of the bend*, so that TE_{01} is coupled to a definite polarization of each spurious mode. The normalized T -functions are then:

$$\begin{aligned} T_{(nm)} &= \sqrt{\frac{\epsilon_n}{\pi}} \frac{J_n(\chi_{(nm)}\rho) \sin n\varphi}{k_{(nm)} J_{n-1}(k_{(nm)})}, \\ T_{[nm]} &= \sqrt{\frac{\epsilon_n}{\pi}} \frac{J_n(\chi_{[nm]}\rho) \cos n\varphi}{(k_{[nm]}^2 - n^2)^{\frac{1}{2}} J_n(k_{[nm]})}, \end{aligned} \quad (42)$$

where

$$\begin{aligned} k_{(nm)} &= \chi_{(nm)}a, & J_n(k_{(nm)}) &= 0, \\ k_{[nm]} &= \chi_{[nm]}a, & J_n'(k_{[nm]}) &= 0, \end{aligned} \quad (43)$$

and

$$\epsilon_n = \begin{cases} 1, & n = 0, \\ 2, & n \neq 0. \end{cases} \quad (44)$$

1.3.1 Coupling Coefficients due to Curvature

We know that to first order the curvature of the guide can couple the TE_{01} mode only to modes of angular index unity. Let us calculate the

coupling between TE_{01} and TM_{1m} . Referring to (31), we have

$$\begin{aligned} \Xi_{(1m)[01]} &= \Xi_{[01](1m)} \\ &= \int_S \xi(\text{grad } T_{(1m)}) \cdot (\text{flux } T_{[01]}) \, dS \\ &= \int_0^{2\pi} \int_0^a \frac{\sqrt{2} J_1(\chi_{[01]}\rho) J_1(\chi_{(1m)}\rho) \cos^2 \varphi}{\pi a b k_{(1m)} J_0(k_{[01]}) J_0(k_{(1m)})} \rho \, d\rho \, d\varphi \\ &= \begin{cases} a/\sqrt{2} k_{[01]} b & \text{if } m = 1, \\ 0 & \text{if } m \neq 1. \end{cases} \end{aligned} \tag{45}$$

Hence the only transverse magnetic mode coupled to TE_{01} by the bend is TM_{11} , and from (41) the forward coupling coefficient is:

$$c_{(11)[01]}^+ = c_{[01](11)}^+ = \beta a / \sqrt{2} \, k_{[01]} b = 0.18454 \beta a / b. \tag{46}$$

To obtain the coupling between TE_{01} and TE_{1m} , we must evaluate two integrals. From (28), the first is

$$\begin{aligned} \xi_{[01][1m]} &= \xi_{[1m][01]} \\ &= \chi_{[01]} \chi_{[1m]} \int_S \xi T_{[01]} T_{[1m]} \, dS \\ &= \int_0^{2\pi} \int_0^a \frac{\sqrt{2} \chi_{[1m]} \rho^2 J_0(\chi_{[01]}\rho) J_1(\chi_{[1m]}\rho) \cos^2 \varphi}{\pi a b (k_{[1m]}^2 - 1)^{\frac{1}{2}} J_0(k_{[01]}) J_1(k_{[1m]})} \, d\rho \, d\varphi \\ &= \frac{\sqrt{2} a}{b} \frac{k_{[1m]} (k_{[01]}^2 + k_{[1m]}^2)}{(k_{[1m]}^2 - 1)^{\frac{1}{2}} (k_{[01]}^2 - k_{[1m]}^2)^2}, \end{aligned} \tag{47}$$

and from (31), the second is

$$\begin{aligned} \Xi_{[01][1m]} &= \Xi_{[1m][01]} \\ &= \int_S \xi(\text{grad } T_{[01]}) \cdot (\text{grad } T_{[1m]}) \, dS \\ &= - \int_0^{2\pi} \int_0^a \frac{\sqrt{2} \chi_{[1m]} \rho^2 J_1(\chi_{[01]}\rho) J_1'(\chi_{[1m]}\rho) \cos^2 \varphi}{\pi a b (k_{[1m]}^2 - 1)^{\frac{1}{2}} J_0(k_{[01]}) J_1(k_{[1m]})} \, d\rho \, d\varphi \\ &= \frac{2\sqrt{2} a}{b} \frac{k_{[01]} k_{[1m]}^2}{(k_{[1m]}^2 - 1)^{\frac{1}{2}} (k_{[01]}^2 - k_{[1m]}^2)^2}. \end{aligned} \tag{48}$$

A short table of numerical values of the above two integrals follows:

m	$\xi_{[01][1m]}$	$\Xi_{[01][1m]}$
1	0.23871 a/b	0.18638 a/b
2	0.32865 a/b	0.31150 a/b
3	0.03682 a/b	0.02751 a/b

Putting these values into the expression (41) for the forward coupling coefficient, we obtain:

$$\begin{aligned} c_{[01][11]}^+ &= \left[\frac{0.09319(\beta a)^2 - 0.84204}{\sqrt{h_{[01]}ah_{[11]}a}} + 0.09319 \sqrt{h_{[01]}ah_{[11]}a} \right] \frac{1}{b}, \\ c_{[01][12]}^+ &= \left[\frac{0.15575(\beta a)^2 - 3.35688}{\sqrt{h_{[01]}ah_{[12]}a}} + 0.15575 \sqrt{h_{[01]}ah_{[12]}a} \right] \frac{1}{b}, \\ c_{[01][13]}^+ &= \left[\frac{0.01376(\beta a)^2 - 0.60216}{\sqrt{h_{[01]}ah_{[13]}a}} + 0.01376 \sqrt{h_{[01]}ah_{[13]}a} \right] \frac{1}{b}. \end{aligned} \quad (49)$$

1.3.2 Coupling Coefficients due to Dielectric

Depending upon the distribution of the dielectric material over the cross section of the guide, the TE_{01} mode may be coupled to any mode except those of the TM_{0m} family. The dielectric coupling coefficients, as given by equations (41), are

$$\begin{aligned} d_{[01](nm)}^+ &= d_{(nm)[01]}^+ = \frac{1}{2} \beta \Delta_{[01](nm)} \sqrt{h_{(nm)}/h_{[01]}}, \\ d_{[01][nm]}^+ &= d_{[nm][01]}^+ = \frac{\beta^2 \Delta_{[01][nm]}}{2 \sqrt{h_{[01]}h_{[nm]}}}. \end{aligned} \quad (50)$$

The Δ 's are obtained from equations (32); thus,

$$\begin{aligned} \Delta_{[01](nm)} &= \int_S \delta(\text{grad } T_{(nm)}) \cdot (\text{flux } T_{[01]}) dS \\ &= \int_0^{2\pi} \int_0^a \frac{\delta(\rho, \varphi) n \sqrt{\epsilon_n} J_n(\chi_{(nm)} \rho) J_1(\chi_{[01]} \rho) \cos n\varphi}{\pi a k_{(nm)} J_{n-1}(k_{(nm)}) J_0(k_{[01]})} d\rho d\varphi, \\ \Delta_{[01][nm]} &= \int_S \delta(\text{grad } T_{[nm]}) \cdot (\text{grad } T_{[01]}) dS \\ &= - \int_0^{2\pi} \int_0^a \frac{\delta(\rho, \varphi) \sqrt{\epsilon_n} \chi_{[nm]} J_n'(\chi_{[nm]} \rho) J_1(\chi_{[01]} \rho) \cos n\varphi}{\pi a (k_{[nm]}^2 - n^2)^{\frac{1}{2}} J_n(k_{[nm]}) J_0(k_{[01]})} \rho d\rho d\varphi. \end{aligned} \quad (51)$$

1.4 Dielectric Losses

To take account of dielectric dissipation we may introduce a complex permittivity,

$$\epsilon = \epsilon' - i\epsilon'' = \epsilon' (1 - i \tan \varphi), \quad (52)$$

where φ is the loss angle of the dielectric and is not, of course, to be confused with the coordinate φ . If we let

$$\epsilon' = \epsilon_0 (1 + \delta'), \quad (53)$$

where δ' is real, then (52) may be written

$$\epsilon = \epsilon_0[1 + \delta' - i(1 + \delta') \tan \varphi] = \epsilon_0(1 + \delta), \quad (54)$$

where

$$\delta = \delta' - i(1 + \delta') \tan \varphi. \quad (55)$$

No changes in the formal analysis result from the fact that δ now has a complex value.

If the compensator is designed (assuming a lossless dielectric) so that the total power coupled from TE_{01} into all spurious modes is small at all points, it is reasonable to assume that the principal effect of dielectric loss on the TE_{01} mode will be seen in the modified phase constant $\beta_{[01]}$ of this mode in the presence of the compensator. If the compensator is made by filling a certain part S_1 of the guide cross section with a medium of constant (complex) permittivity, and the rest of the cross section with air, then from (32) and (40) the modified phase constant of the TE_{01} mode is

$$\beta_{[01]} = h_{[01]} + \frac{\beta^2 \delta}{2h_{[01]}} \int_{S_1} (\text{grad } T_{[01]})^2 dS, \quad (56)$$

where δ is given by (55). Since δ is complex, the attenuation constant is

$$\alpha_{[01]} = -\text{Im } \beta_{[01]} = \frac{\beta^2(1 + \delta') \tan \varphi}{2h_{[01]}} \int_{S_1} (\text{grad } T_{[01]})^2 dS, \quad (57)$$

where the integration may be carried out as soon as the area S_1 is specified.

The approximation (57) for the attenuation constant due to dielectric losses has a simple physical interpretation. It corresponds to the power which would be dissipated in a medium of conductivity $\omega\epsilon''$ if the electric field existing in the medium were the same as the field of the TE_{01} mode in a straight, empty guide. This is probably a very good approximation if δ' is small, as it will be for the foam dielectrics from which compensators are most likely to be made.

It is doubtful that (57) furnishes a good approximation to the dielectric loss when the permittivity of the compensator is high (δ not small compared to unity). If the permittivity is high the cross section of the dielectric member will be small, but the field perturbation may be large in the immediate neighborhood of the dielectric. The series which represent the fields in terms of the normal modes of the empty guide may converge slowly; in other words, when using the telegraphist's equations one must consider the coupling between TE_{01} and a large number of

other modes, none of which appears by itself to be very strongly coupled to TE_{01} .

Of course if we had a single normal mode of the *compensated* guide, with a field pattern independent of distance along the guide, it might well be possible to calculate the field distribution and the dielectric losses approximately, without reference to the telegraphist's equations and regardless of the permittivity of the dielectric. However, we do not have a single normal mode of the compensated guide, but rather a mixture of modes. The field pattern varies along the guide as the modes phase in and out; and it is not easy to conclude from this picture what the actual dielectric losses will be.

Finally it should be remembered that we have said nothing about the possible effect of a dielectric compensator on eddy current losses in the waveguide walls. If one attempted to use a compensator of small physical size and correspondingly high permittivity, the resulting perturbation of the electric field might very well increase the eddy current losses in the wall adjacent to the compensator. On the other hand, the increase would probably be negligible for a compensator made out of a foam dielectric.

II. APPLICATION

2.1 *Properties of Uniformly Coupled Transmission Lines*

We shall now apply the preceding theory to the calculation of TE_{01} mode coupling in gentle bends. To describe propagation in a curved waveguide in terms of the modes of a straight guide requires, in general, the solution of the infinite set of equations (37); but we can give an adequate approximate treatment by considering just two modes at a time, one mode of each pair always being TE_{01} . Furthermore we need consider only the forward waves, since the relative power coupled from the forward waves into the backward waves is quite small.

The differential equations representing the forward waves on two uniformly coupled transmission lines are:

$$\begin{aligned}\frac{da_0}{dz} + \gamma_0 a_0 + i\kappa a_1 &= 0, \\ i\kappa a_0 + \frac{da_1}{dz} + \gamma_1 a_1 &= 0.\end{aligned}\tag{58}$$

In these equations $a_0(z)$ and $a_1(z)$ are the amplitudes of the forward traveling waves, normalized so that $|a_0|^2$ and $|a_1|^2$ represent power flow directly. We may think of the subscript 0 as always referring to the TE_{01}

mode. The complex constants γ_0 and γ_1 may be regarded as (modified) propagation constants; note that because of the coupling they are not necessarily equal to the propagation constants of the uncoupled modes. The coupling coefficient is denoted by κ ; it will be real if the coupling mechanism is lossless, but is not required to be so in the general mathematical solution.

We are interested in the case in which line 0 contains unit power at $z = 0$, and line 1 contains no power. Subject to the initial conditions

$$a_0(0) = 1, \quad a_1(0) = 0, \quad (59)$$

the solution of (58) is

$$\begin{aligned} a_0(z) &= \left[\frac{1}{2} + \frac{(\gamma_0 - \gamma_1)}{2r} \right] e^{m_2 z} + \left[\frac{1}{2} - \frac{(\gamma_0 - \gamma_1)}{2r} \right] e^{m_1 z}, \\ a_1(z) &= \frac{i\kappa}{r} [e^{m_2 z} - e^{m_1 z}], \end{aligned} \quad (60)$$

where

$$r = \sqrt{(\gamma_0 - \gamma_1)^2 - 4\kappa^2}, \quad (61)$$

and

$$\begin{aligned} m_1 &= \frac{1}{2}[-(\gamma_0 + \gamma_1) + r], \\ m_2 &= \frac{1}{2}[-(\gamma_0 + \gamma_1) - r]. \end{aligned} \quad (62)$$

For the case of two propagating modes without loss, κ is real and we may write

$$\gamma_0 = i\beta_0, \quad \gamma_1 = i\beta_1, \quad (63)$$

so that

$$r = i\sqrt{(\beta_0 - \beta_1)^2 + 4\kappa^2}. \quad (64)$$

A straightforward calculation now gives the power in each line at any point:

$$\begin{aligned} P_0 &= |a_0(z)|^2 = 1 - \frac{4\kappa^2}{(\beta_0 - \beta_1)^2 + 4\kappa^2} \sin^2 \frac{1}{2}[(\beta_0 - \beta_1)^2 + 4\kappa^2]^{\frac{1}{2}} z, \\ P_1 &= |a_1(z)|^2 = \frac{4\kappa^2}{(\beta_0 - \beta_1)^2 + 4\kappa^2} \sin^2 \frac{1}{2}[(\beta_0 - \beta_1)^2 + 4\kappa^2]^{\frac{1}{2}} z. \end{aligned} \quad (65)$$

Hence the maximum power transferred from line 0 to line 1 is

$$(P_1)_{\max} = \frac{4\kappa^2}{(\beta_0 - \beta_1)^2 + 4\kappa^2} = \frac{[2\kappa/(\beta_0 - \beta_1)]^2}{1 + [2\kappa/(\beta_0 - \beta_1)]^2}, \quad (66)$$

and the points of maximum power transfer are

$$z = \frac{(2n + 1)\pi}{\sqrt{(\beta_0 - \beta_1)^2 + 4\kappa^2}}, \quad (67)$$

where n is any integer.

As is well known, complete power transfer from line 0 to line 1 is possible if and only if the (modified) phase constants β_0 and β_1 are equal. In general the maximum power transferred to line 1 depends on the ratio of the coupling coefficient to the difference in phase constants, and if this ratio is small, then

$$(P_1)_{\max} \approx 4\kappa^2/(\beta_0 - \beta_1)^2. \quad (68)$$

If the difference in phase constants is not large enough to prevent undesirable power loss from line 0, an alternative possibility, as Miller⁸ has shown, is to increase the attenuation constant of line 1 while leaving the attenuation constant of line 0 as nearly unchanged as possible. We can get an idea of the required attenuation difference from the following approximate treatment.

Let

$$\begin{aligned} \gamma_0 &= \alpha_0 + i\beta_0, \\ \gamma_1 &= \alpha_1 + i\beta_1, \end{aligned} \quad (69)$$

and assume that

$$|2\kappa/(\gamma_0 - \gamma_1)|^2 \ll 1. \quad (70)$$

Then

$$r \approx \gamma_0 - \gamma_1 - \frac{2\kappa^2}{(\gamma_0 - \gamma_1)}, \quad (71)$$

and

$$\begin{aligned} m_1 &\approx -\gamma_1 - \frac{\kappa^2}{\gamma_0 - \gamma_1}, \\ m_2 &\approx -\gamma_0 + \frac{\kappa^2}{\gamma_0 - \gamma_1}. \end{aligned} \quad (72)$$

We are interested in the power in line 0. From the first of equations (60), we have

$$a_0(z) \approx \left[1 + \frac{\kappa^2}{(\gamma_0 - \gamma_1)^2}\right] e^{m_2 z} - \frac{\kappa^2}{(\gamma_0 - \gamma_1)^2} e^{m_1 z}. \quad (73)$$

Let us consider the case in which line 1 has a much higher attenuation constant than line 0; that is,

$$\alpha_1 \gg \alpha_0. \quad (74)$$

The second term on the right side of (73) is provided with a small coefficient, and also its exponential factor decays much faster than the exponential in the first term. The second term, therefore, rapidly becomes negligible as z increases, and we may write,

$$a_0(z) \approx \left[1 + \frac{\kappa^2}{(\gamma_0 - \gamma_1)^2} \right] e^{\kappa^2 z / (\gamma_0 - \gamma_1)} e^{-\gamma_0 z}. \quad (75)$$

If the attenuation constant of line 0 is not modified* by the presence of the coupled lossy line 1, then in the absence of line 1 the amplitude of $a_0(z)$ would be $e^{-\alpha_0 z}$, and the factor by which the amplitude is reduced owing to the presence of line 1 is

$$\left| 1 + \frac{\kappa^2}{(\gamma_0 - \gamma_1)^2} \right| e^{\kappa^2 z / (\gamma_0 - \gamma_1)}. \quad (76)$$

The first factor on the right is very nearly unity, but not less than unity if κ is real (lossless coupling mechanism) and

$$(\alpha_1 - \alpha_0)^2 \geq (\beta_1 - \beta_0)^2. \quad (77)$$

Hence the factor by which the amplitude is multiplied is not less than

$$\left| \exp \frac{\kappa^2 z}{\gamma_0 - \gamma_1} \right| = \exp \frac{(\alpha_0 - \alpha_1) \kappa^2 z}{(\alpha_0 - \alpha_1)^2 + (\beta_0 - \beta_1)^2}, \quad (78)$$

assuming that κ is real. If the amplitude of the wave on line 0 is not to be down by more than N nepers, after a distance z , from what it would have been in the absence of the coupled line, it suffices to have

$$\frac{(\alpha_1 - \alpha_0) \kappa^2 z}{(\alpha_1 - \alpha_0)^2 + (\beta_1 - \beta_0)^2} = N, \quad (79)$$

or

$$\alpha_1 - \alpha_0 = \frac{1}{2} \left[(\kappa^2 z / N) + \sqrt{(\kappa^2 z / N)^2 - 4(\beta_1 - \beta_0)^2} \right]. \quad (80)$$

2.2 TE₀₁-TM₁₁ Coupling in Plain and Compensated Bends

In Jouguet's¹ and Rice's² analysis of propagation in a curved waveguide, the curvature is treated as a perturbation and the field com-

* The value of α_0 may very well be modified by the coupling; but if it is this can easily be taken into account when computing the over-all change in $|a_0(z)|$ due to the presence of line 1.

ponents are developed in powers of the small parameter a/b , but the field perturbations are not expressed in terms of the modes of the straight waveguide. We shall now consider propagation in plain and compensated bends from the coupled-mode viewpoint.

Denote the coefficient of coupling between the TE_{01} mode and the TM_{nm} mode or the TE_{nm} mode by

$$\begin{aligned} \kappa_{(nm)} &= c_{(nm)} + d_{(nm)}, \\ \kappa_{[nm]} &= c_{[nm]} + d_{[nm]}, \end{aligned} \quad (81)$$

respectively, where as usual we indicate TM modes by enclosing the subscripts in *parentheses* and TE modes by enclosing the subscripts in *brackets*. In Part I the coupling coefficients are written with two pairs of subscripts, since in general they may refer to any two modes, but here one pair of subscripts would always be [01] and will be omitted. The coefficient c represents that part of the coupling (if any) which is due to the curvature of the guide, and d represents the coupling (if any) which is due to the inhomogeneity of the dielectric. We assume that the dielectric loading, if present, is symmetric with respect to the plane of the bend, so that coupling to only one polarization of each mode need be considered.

The phase constants of the TM_{nm} and TE_{nm} modes in a straight, empty guide are, respectively,

$$h_{(nm)} = \sqrt{\beta^2 - \chi_{(nm)}^2}, \quad h_{[nm]} = \sqrt{\beta^2 - \chi_{[nm]}^2}, \quad (82)$$

where

$$\chi_{(nm)} = k_{(nm)}/a, \quad \chi_{[nm]} = k_{[nm]}/a, \quad (83)$$

and

$$J_n(k_{(nm)}) = 0, \quad J_n'(k_{[nm]}) = 0. \quad (84)$$

Also

$$\beta = 2\pi/\lambda, \quad (85)$$

where λ is the free-space wavelength.

As noted in the preceding section, the presence of coupling may cause the modified phase constants $\beta_{(nm)}$ and $\beta_{[nm]}$ of the coupled modes to differ slightly from the unperturbed phase constants $h_{(nm)}$ and $h_{[nm]}$. In a plain bend (curvature coupling only), the β 's are equal to the h 's, and in most cases the effect of a small amount of dielectric coupling on the phase constants may be neglected. Exact values of $\beta_{[nm]}$ and $\beta_{(nm)}$ may be obtained if necessary from (40).

The coupling coefficient between the TE_{01} and TM_{11} modes in a plain bend is given by equation (46) as

$$c_{(11)} = \beta a / \sqrt{2} k_{[01]} b = 0.18454 \beta a / b = 1.1595a / \lambda b. \quad (86)$$

The (smallest) critical distance for maximum power transfer is given by (67) with $\beta_0 = \beta_1$, namely

$$z_{c_0} = \frac{\pi}{2c_{(11)}} = \frac{k_{[01]} b \lambda}{2\sqrt{2}a} = \frac{1.3547b\lambda}{a}, \quad (87)$$

and the critical angle ϑ_{c_0} is

$$\vartheta_{c_0} = z_{c_0} / b = 1.3547 \lambda / a \text{ radians} = 77.62 \lambda / a \text{ degrees}. \quad (88)$$

This expression agrees, as it must, with that obtained by Jouguet and Rice. (We write ϑ_{c_0} for the uncompensated bend in order to reserve ϑ_c for a bend with dielectric loading.)

It should be pointed out that $c_{(11)}$ is not necessarily the largest of the coupling coefficients due to curvature. In a guide sufficiently far above cutoff, it appears from (49) that $c_{[11]}$ is approximately equal to $c_{(11)}$, and $c_{[12]}$ is one and two-thirds times as large as $c_{(11)}$. If two transmission lines are coupled over a distance z which is small compared to the distance required for maximum power transfer, then by (65) the relative power transferred to line 1 is

$$P_1(z) \approx \kappa^2 z^2, \quad (89)$$

which is proportional to the square of the coupling coefficient. It follows that for a sufficiently small bending angle the largest amount of power will go into the mode which has the largest coupling coefficient to TE_{01} (in the above example, TE_{12}). Each coupling coefficient, however, is proportional to $1/b$, and can be made as small as desired by increasing the radius of curvature of the bend. Since the phase constants are unaltered to first approximation by the curvature, the maximum power transferred tends to zero with $1/b^2$ for every mode whose unperturbed phase constant differs from $h_{[01]}$. The only mode with finite power transfer for a finite bending angle with an arbitrarily large bending radius is TM_{11} , since $h_{(11)} = h_{[01]}$. For the present we shall assume a bending radius so large that power transfer to modes other than TM_{11} is negligible. Complete power transfer from TE_{01} to TM_{11} will then take place in a plain bend at odd multiples of the critical bending angle ϑ_{c_0} .

We now consider a dielectric-loaded bend in which the permittivity ϵ is a function of the transverse coordinates (ρ, φ) , but does not vary

from one cross section to another. The permeability μ is assumed to be constant. Thus we shall write, as in (25),

$$\begin{aligned}\mu &= \mu_0, \\ \epsilon &= \epsilon_0[1 + \delta(\rho, \varphi)],\end{aligned}\tag{90}$$

and assume that

$$\frac{1}{S} \int_S |\delta| dS \ll 1,\tag{91}$$

where S is the cross-sectional area of the guide. The inequality (91) implies either that ϵ does not vary much over the cross section, or that it varies extensively over only a small part of the cross section. The first alternative corresponds to a dielectric whose relative permittivity differs but little from unity, and is the most likely case in practice. If, on the other hand, δ is large in a small region (a thin sliver of high-permittivity material), the dielectric coupling coefficients will not diminish very rapidly with increasing mode number. The TE_{01} mode will be appreciably coupled to a large number of modes, and it may not be safe to assume that the total power converted into spurious modes is small just because the conversion into any given mode is small. We shall not try to decide here what maximum value of $|\delta|$ is practicable.

If the distribution of dielectric is symmetric with respect to the plane of the bend, the dielectric couples the TE_{01} mode to a definite polarization of each spurious mode, and in particular to the same polarization of the TM_{11} mode that is coupled by curvature alone.* The dielectric coupling coefficient between TE_{01} and TM_{11} is, from (50) and (51),

$$d_{(11)} = \frac{\beta}{\sqrt{2\pi}ak_{[01]}J_0^2(k_{[01]})} \int_S \delta(\rho, \varphi) J_1^2(\chi_{[01]}\rho) \frac{\cos \varphi}{\rho} dS.\tag{92}$$

It is obvious that the *decoupling condition*, namely:

$$\kappa_{(11)} \equiv c_{(11)} + d_{(11)} = 0,\tag{93}$$

may be satisfied by an infinite number of different distributions of permittivity. Ingenuity is required, however, to find a configuration which is easy to fabricate and which does not couple the TE_{01} mode too strongly to any other mode in the guide. The spurious mode problem is quite serious when the diameter of the guide (in wavelengths) is so large that

* A dielectric insert which is not symmetric with respect to the plane of the bend will couple TE_{01} to the other polarization of TM_{11} , with potentially complete power transfer to this mode on account of the equality of phase velocities. A small accidental lack of symmetry should not lead to serious mode conversion in a bend of moderate angle.

many modes have phase velocities close to TE_{01} . In the next section we shall compute the coupling to various modes for a number of cases.

Mention may be made here of the effect of imperfect decoupling. If we could satisfy the decoupling condition (93) exactly, then to first order there would be no transfer of power from TE_{01} to TM_{11} in a bend of any angle. In practice we cannot expect to satisfy (93) exactly, on account of uncertainties in the permittivity and dimensions of the compensator. If the coupling coefficients are constant along the bend, the effect of reducing $\kappa_{(11)}$ by making $d_{(11)}$ nearly equal and opposite to $c_{(11)}$ is to increase the distance required for maximum power transfer to TM_{11} to take place. In the simple case where $\beta_{(11)} = \beta_{[01]}$, as, for example, when the compensator is made of a bent half-cylinder of dielectric, the critical angle for a fixed bend radius is inversely proportional to $\kappa_{(11)}$. The critical angle ϑ_c for an imperfectly compensated bend, in terms of the critical angle ϑ_{c_0} for an uncompensated bend, is

$$\vartheta_c = \left| \frac{c_{(11)}}{c_{(11)} + d_{(11)}} \right| \vartheta_{c_0}. \quad (94)$$

If $d_{(11)}$ can be made equal and opposite to $c_{(11)}$ within 1 per cent, say, then $\vartheta_c = 100 \vartheta_{c_0}$. The power transferred in a bend of angle ϑ is simply proportional to $\sin^2 (\vartheta/\vartheta_c)$.

2.3 Various Compensator Designs

From now on we shall consider a compensated bend in which the TE_{01} and TM_{11} modes are completely decoupled, and we shall investigate the coupling between TE_{01} and spurious modes. (By "spurious mode" we mean any mode of the straight round guide except TE_{01} or TM_{11} .) We shall assume that the power in all spurious modes is everywhere low compared to the power level of TE_{01} . To first order, therefore, we may compute the power abstracted from TE_{01} by mode conversion by assuming that the TE_{01} mode crosstalks into each spurious mode independently. In practice the crosstalk is greatest into modes whose phase velocities are nearest to the phase velocity of TE_{01} ; and it seems more than adequate, at least for foam dielectric compensators, to consider about a dozen modes.

From (68), the maximum relative power (assumed small) which cross-talks from TE_{01} into a given spurious mode is

$$(P_1)_{\max} \approx [2\kappa/(\beta_0 - \beta_1)]^2, \quad (95)$$

where κ is the total coupling coefficient. For all spurious modes we assume that the difference in phase constants is not much changed by

the dielectric loading; this assumption obviates the somewhat laborious calculation of β_1 for each spurious mode from (40). Thus,

$$\beta_0 - \beta_1 \approx h_0 - h_1, \quad (96)$$

and a good estimate of the maximum power which crosstalks into any spurious mode is

$$(P_1)_{\max} \approx [2(c + d)/(h_0 - h_1)]^2. \quad (97)$$

If the form and dimensions of a bend compensator are fixed, and the TE₀₁—TM₁₁ decoupling condition is assumed to be met by adjusting the permittivity, it turns out that the maximum power which crosstalks into a given mode is proportional to $(a/b)^2$. It is thus easy to calculate the bending radius which makes $(P_1)_{\max}$ for any given mode equal to a (small) preassigned value. The total power abstracted from the TE₀₁ mode by mode conversion will be a complicated, fluctuating function of distance along the bend, or of frequency at a fixed distance, because of the different critical distances for maximum power transfer into the different spurious modes, but we can get an idea of the minimum tolerable bending radius by considering just the crosstalk into the one or two most troublesome modes.

It seems likely that with present-day dielectrics at millimeter wavelengths, dielectric losses in a compensated bend will be comparable to mode conversion losses. For this reason an estimate of dielectric losses is given in connection with each type of compensator design discussed below.

2.3.1 *The Geometrical Optics Solution*

An obvious way to design a bend compensator on paper is to load the bend with a medium of continuously varying permittivity for which*

$$\delta = -2(\rho/b) \cos \varphi. \quad (98)$$

This may be called the geometrical optics solution of the bend problem, since to first order it equalizes the optical length of all circular paths which are coaxial with the curved center line of the waveguide. Physically the required variation of permittivity is rather simple; the permittivity at each point depends only on the distance of the point from a line through the center of curvature perpendicular to the plane of the bend. An attempt to indicate this variation by shading has been made

* In order that the relative permittivity of the medium be nowhere less than unity, the constant ϵ_0 in the expression $\epsilon = \epsilon_0(1 + \delta)$ must be greater than the permittivity of free space; but this does not affect the analysis in any way.

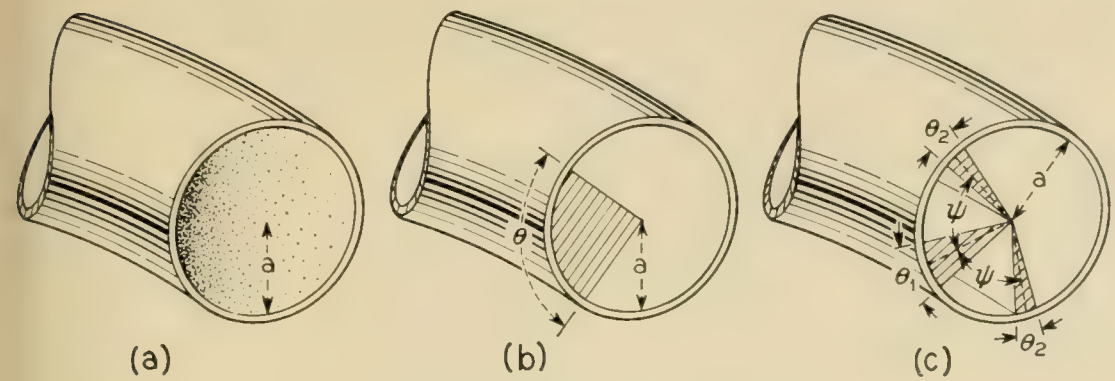


Fig. 2 — Various bend compensators.

in Fig. 2(a). One could approximate the continuous variation with a laminated structure consisting of a number of layers of different permittivities, the permittivity varying slightly from one layer to the next.

Although the geometrical optics approach provides a very good theoretical solution to the bend problem, it does not lead to a perfect compensator. It is shown at the end of this section that a perfect compensator, in the sense of one which does not couple TE_{01} to any other mode at any frequency, does not exist. The geometrical optics compensator couples the same modes to TE_{01} that are coupled by the curvature of the bend itself, namely TM_{11} and the TE_{1m} family; but the coupling coefficients of the various modes are not in exactly the same ratios. Thus if δ is given by (98), the net coupling between TE_{01} and TM_{11} in the compensated bend is zero, but there will be a small residual coupling between TE_{01} and each of the TE_{1m} modes.

The curvature coupling coefficients are given by (49). Table I contains numerical values which have been worked out for $\beta a = 12.930$ and $\beta a = 29.554$, corresponding respectively to guide diameters of $\frac{7}{8}$ inch and 2 inches at an operating wavelength of 5.4 mm.

Dielectric coupling coefficients can be worked out without much dif-

TABLE I — COUPLING COEFFICIENTS AND POWER TRANSFER
IN GEOMETRICAL OPTICS COMPENSATOR

Mode	$\beta a = 12.930$			$\beta a = 29.554$		
	c	d	$2 \frac{(c + d)}{(h_0 - h_1)}$	c	d	$2 \frac{(c + d)}{(h_0 - h_1)}$
TM_{11}	$2.386/b$	$-2.386/b$	—	$5.454/b$	$-5.454/b$	—
TE_{11}	$2.344/b$	$-2.479/b$	$0.600a/b$	$5.480/b$	$-5.537/b$	$0.597a/b$
TE_{12}	$3.759/b$	$-4.318/b$	$-1.962a/b$	$9.092/b$	$-9.322/b$	$-1.954a/b$
TE_{13}	$0.306/b$	$-0.420/b$	$-0.087a/b$	$0.793/b$	$-0.834/b$	$-0.083a/b$

ficiency from (50), (51), and (98). The first few coefficients are:

$$\begin{aligned} d_{(11)} &= -0.18454\beta a/b, \\ d_{[11]} &= -\frac{0.18638(\beta a)^2/b}{\sqrt{h_{[01]}ah_{[11]}a}}, \\ d_{[12]} &= -\frac{0.31150(\beta a)^2/b}{\sqrt{h_{[01]}ah_{[12]}a}}, \\ d_{[13]} &= -\frac{0.02751(\beta a)^2/b}{\sqrt{h_{[01]}ah_{[13]}a}}. \end{aligned} \tag{99}$$

Some numerical values are recorded in Table I.

The phase constant of any mode in the compensated bend is equal to the phase constant of the same mode in a guide filled with material of constant permittivity ϵ_0 . We may therefore set $\beta_{[01]} - \beta_{[1m]}$ equal to $h_{[01]} - h_{[1m]}$. Strictly speaking, λ is then the wavelength of a free wave in a medium of permittivity ϵ_0 ; but ϵ_0 differs little from the permittivity of vacuum if the compensator is made from a foam dielectric. The ratio of total coupling coefficient to difference in phase constants, namely $2(c + d)/(h_0 - h_1)$, which determines the maximum power transfer by (97), is given in Table I for the $\frac{7}{8}$ -inch and 2-inch guides at $\lambda = 5.4$ mm.

For large βa the leading terms in $c_{[1m]}$ and $d_{[1m]}$, which are proportional to βa , cancel each other, and $c_{[1m]} + d_{[1m]}$ decreases like $1/\beta a$. The difference in phase constants, $h_{[01]} - h_{[1m]}$, is also proportional to $1/\beta a$ for large βa . Hence the ratio of coupling coefficient to difference in phase constants approaches a finite limit as βa approaches infinity; to three decimal places the limiting values are the same as those given in Table I for $\beta a = 29.554$.

If we choose a sufficiently large value of a/b , the maximum power transferred to a given mode may be made to approach any preassigned small value as λ/a approaches zero. *This is a special property of the geometrical optics compensator.* For other compensator designs $c + d$ will be of the order of βa while $h_0 - h_1$ will be of the order of $1/\beta a$, so that $2(c + d)/(h_0 - h_1)$ will tend to infinity like $(\beta a)^2$. Hence in the short-wavelength limit the bend output will be a jumble of modes all at comparable power levels. Practically, one must expect the same end result with a geometrical optics compensator, because of the impossibility of meeting exact mathematical specifications; but a carefully-designed geometrical optics compensator should work satisfactorily up to a considerably higher frequency than any other type.

To get an idea of tolerable bending radii with a geometrical optics compensator, we shall calculate the radius at which the maximum mode conversion loss from TE_{01} into TE_{12} (the worst spurious mode) is 0.1 db. Setting $P_1 = 0.02276$, we find from (68) that the minimum bending radius for a $\frac{7}{8}$ -inch guide is 5.69 inches, and for a 2-inch guide, 12.95 inches, both at a wavelength of 5.4 mm. It is worth noting that if $|d_{[12]}|$ is increased by 5 per cent of its theoretical value, the minimum bending radius becomes 7.89 inches for the $\frac{7}{8}$ -inch guide and 39.2 inches for the 2-inch guide. (This assumes that the TM_{11} mode is still properly decoupled and that TE_{12} is still the worst spurious mode.)

Dielectric losses are likely to be a serious problem in a geometrical optics compensator, inasmuch as the whole volume of the bent guide has to be filled with dielectric. Relatively large values of δ are required to negotiate bends as sharp as those just discussed. For example, if $b = 13a$, in a practical case δ might range from 0.058 at the inner surface of the bend to 0.250 at the center of the guide to 0.442 at the outer surface (referred to ϵ_0 as the permittivity of free space). The loss tangent of present-day dielectrics in this range is approximately 2×10^{-4} . A large (i.e., far above cutoff) waveguide filled with a dielectric of relative permittivity 1.25 and loss tangent 2×10^{-4} will show a dielectric loss of about 1.13 db/meter or 0.34 db/ft at 5.4 mm. The dielectric loss in a 90° bend with a bending radius of 1 foot would be about 0.54 db, and for other bends the loss would be directly proportional to the length of the bend, and to the loss tangent if different from 2×10^{-4} . It is true that loss tangents as low as 5×10^{-5} may be obtained with lower values of permittivity, say $\delta = 0.033$; but with such a small δ the bend radius must be proportionately larger, and the total dielectric losses in a bend of given angle would be larger than with a higher permittivity material.

We proceed now to demonstrate the assertion made earlier that a perfect bend compensator does not exist. More precisely, we shall show that it is impossible to compensate the bend with an isotropic medium whose permittivity and permeability are everywhere finite, but otherwise arbitrary, in such a way that there is no conversion from TE_{01} to any other mode at any point at any frequency.

If there is no mode conversion at any point of the bend then the fields at all points must be those of the TE_{01} mode, referred to the bent cylindrical coordinate system described at the beginning of Section 1.1. In other words, we have prescribed the electromagnetic field and are asking whether it is possible to choose the permeability and permittivity so that the given field will satisfy Maxwell's equations. Usually the

answer to this question will be "No." The Maxwell equations,

$$\begin{aligned}\vec{\nabla} \times \vec{E} &= -i\omega\mu\vec{H}, \\ \vec{\nabla} \times \vec{H} &= i\omega\epsilon\vec{E},\end{aligned}\tag{100}$$

are equivalent to six scalar equations, and if the components of E and $\vec{H}$ are prescribed, one cannot in general satisfy all these equations by merely adjusting the two scalar functions μ and ϵ .

It is particularly easy to see the difficulty for the TE_{01} mode in a curved guide. Recall that the TE_{01} mode fields are independent of the coordinate φ , and that the only non-vanishing field components are $E_\varphi(\rho, z)$, $H_\rho(\rho, z)$, and $H_z(\rho, z)$. The fourth of equations (5) is:

$$\frac{1}{\rho[1 + (\rho/b) \cos \varphi]} \left[\frac{\partial}{\partial \varphi} ([1 + (\rho/b) \cos \varphi] H_z) - \frac{\partial}{\partial z} (\rho H_\varphi) \right] = i\omega\epsilon E_\rho.\tag{101}$$

Since $E_\rho = 0$, the right side of the equation is zero for any finite value of ϵ at any finite frequency, and the whole equation reduces to

$$-\frac{H_z \sin \varphi}{b + \rho \cos \varphi} = 0,\tag{102}$$

which can be true for all values of ρ and φ only if the radius of curvature of the bend is infinite. Hence a perfect compensator cannot be designed with *any* value of ϵ .

The practical importance of this result does not appear to be great, since theoretically the geometrical optics solution would provide an extraordinarily good compensator. Until one has a dielectric whose permittivity is continuously variable and precisely controllable, and whose loss tangent is very low, even this solution is of only academic interest.

2.3.2 The Single-Sector Compensator

In practice the simplest way to compensate a bend is to fill part of the cross section of the guide with homogeneous dielectric material of relative permittivity $(1 + \delta)$, and leave the remainder empty. In many cases a suitable shape for the cross section of the dielectric is a sector of a circle, inserted on the side of the guide nearest the center of curvature of the bend, and symmetrically placed with respect to the plane of the bend. Such a sector, of total angle θ , is shown in Fig. 2(b). We shall now discuss the properties of a single-sector compensator.

For future reference, the modified phase constants of the TE_{01} and TM_{11} modes may be calculated from (40) as:

$$\beta_{(11)} = h_{(11)} + \frac{\delta}{4\pi h_{(11)}} [(\theta - 0.29646 \sin \theta)\beta^2 - (0.70354 \sin \theta)\chi_{(11)}^2], \quad (103)$$

$$\beta_{[01]} = h_{[01]} + \frac{\beta^2 \delta \theta}{4\pi h_{[01]}}.$$

The dielectric coupling coefficients for the single-sector compensator are given by (50) and (51), provided that some of the Bessel functions are integrated by Simpson's rule. In particular, the dielectric coupling coefficient between TE_{01} and TM_{11} is

$$d_{(11)} = -0.12066 \beta \delta \sin \theta/2. \quad (104)$$

Substituting (86) and (104) into (93), we obtain *the decoupling condition for a single-sector compensator*, namely

$$\delta = \frac{1.5295}{\sin \theta/2} \frac{a}{b}. \quad (105)$$

No circular magnetic modes (TM_{0m}) and no higher circular electric modes (TE_{02} , TE_{03} , etc.) are coupled to TE_{01} by the single-sector compensator. We also observe that the coupling coefficients of the TE_{1m} and TM_{1m} modes do not depend upon the sector angle θ so long as δ and θ are related by (105). This is because these modes have the same angular dependence as the TM_{11} mode, which we are trying to compensate.

The most troublesome spurious modes are those whose phase velocities are closest to TE_{01} , namely TE_{21} and TE_{31} . The coupling coefficients and power transfer ratios for these two modes in $\frac{7}{8}$ -inch and 2-inch guide are given in Table II. Either of these modes could be decoupled by proper choice of the sector angle, provided we had a uniform, low-dissipation dielectric with the required value of δ , but it is impossible to decouple both modes at once with a single sector. As a compromise, we might adopt the sector angle which equalizes the maximum power transferred to TE_{21} and TE_{31} (the distances for maximum power transfer are of course not the same for the two modes). This angle is approximately 144° .

Suppose we wish to employ a 144° sector in a $\frac{7}{8}$ -inch guide at 5.4 mm. If the criterion is that TE_{01} is to lose a maximum of 0.1 db each to TE_{21} and TE_{31} , so that the total mode conversion losses are of the order of 0.2

TABLE II — COUPLING COEFFICIENTS AND POWER TRANSFER TO TE₂₁ AND TE₃₁ MODES DUE TO SINGLE-SECTOR COMPENSATOR

Mode	$\beta a = 12.930$		$\beta a = 29.554$	
	d	$2d/(h_0 - h_1)$	d	$2d/(h_0 - h_1)$
TE ₂₁	$\frac{1.1152}{b} \frac{\sin \theta}{\sin \theta/2}$	$-10.38 \frac{a}{b} \frac{\sin \theta}{\sin \theta/2}$	$\frac{2.4727}{b} \frac{\sin \theta}{\sin \theta/2}$	$-54.23 \frac{a}{b} \frac{\sin \theta}{\sin \theta/2}$
TE ₃₁	$-\frac{0.6579}{b} \frac{\sin 3\theta/2}{\sin \theta/2}$	$-10.89 \frac{a}{b} \frac{\sin 3\theta/2}{\sin \theta/2}$	$-\frac{1.4426}{b} \frac{\sin 3\theta/2}{\sin \theta/2}$	$-56.91 \frac{a}{b} \frac{\sin 3\theta/2}{\sin \theta/2}$

db, then the bending radius can be 19.5 inches. The corresponding value of δ is 0.036.

If we try to use a 144° sector in a 2-inch guide at 5.4 mm, with the same mode conversion criterion as before, the bending radius must be so large that no currently available dielectric has a small enough value of δ to satisfy the decoupling condition. We are therefore forced to use a smaller sector angle. With a sector of small angle, TE₃₁ is the worst spurious mode. It turns out that if $\delta = 0.033$ and if TE₀₁ is not to lose more than 0.1 db by conversion into TE₃₁, the minimum bending radius is 1131 inches or 94.3 feet, and the corresponding sector angle is 4.70°.

An approximate formula for the attenuation constant due to dielectric losses in a single-sector compensator is given by (57), provided that δ is small. The result is

$$\alpha_d = \frac{\beta^2(1 + \delta) \tan \varphi}{2h_{[01]}} \frac{\theta^\circ}{360} \text{ nepers/meter,} \tag{106}$$

where $\tan \varphi$ is the loss tangent of the dielectric, and θ° is the sector angle in degrees. As numerical examples, we find that the dielectric loss at 5.4 mm in a $\frac{7}{8}$ -inch guide compensated with a 144° sector having $\delta = 0.036$, $\tan \varphi = 5 \times 10^{-5}$, amounts to 0.085 db in a 90° bend of radius 19.5 inches. In a 2-inch guide compensated with a 4.70° sector ($\delta = 0.033$, $\tan \varphi = 5 \times 10^{-5}$), the dielectric loss is 0.155 db in a 90° bend of radius 94.3 feet.

2.3.3 The Three-Sector Compensator

Although the single-sector compensator should work well in a guide which will propagate only 40 to 50 modes, it does not look so promising for a 200- to 300-mode guide, chiefly because of the unavoidable crosstalk into TE₂₁ and/or TE₃₁. We are therefore led to consider the design of a

three-sector compensator which will not couple either TE_{21} or TE_{31} to TE_{01} .

A three-sector compensator is shown schematically in Fig. 2(c). The angle of the center sector is called θ_1 and the angle of each outer sector θ_2 . Each outer sector makes an angle ψ , measured between center planes, with the center sector.

The condition for decoupling TE_{01} from TM_{11} with the three-sector compensator is

$$\delta = \frac{1.5295}{(2 \cos \psi \sin \theta_2/2 + \sin \theta_1/2)} \frac{a}{b}. \quad (107)$$

If δ and a/b are given, two additional conditions may be imposed upon θ_1 , θ_2 , and ψ . The conditions are taken to be:

$$\begin{aligned} 2 \cos 2\psi \sin \theta_2 + \sin \theta_1 &= 0, \\ 2 \cos 3\psi \sin 3\theta_2/2 + \sin 3\theta_1/2 &= 0. \end{aligned} \quad (108)$$

If equations (108) are satisfied, the compensator does not couple TE_{01} to any modes of the TE_{2m} , TE_{3m} , TM_{2m} , or TM_{3m} families.

To design a three-sector compensator with given values of a/b and δ , one can solve (107) and (108) numerically for θ_1 , θ_2 , and ψ . However, if θ_1 is small ($\leq 20^\circ$, for example), a simpler design procedure may be used. The following equations are approximately true:

$$\begin{aligned} \theta_1 &= \frac{126.8a}{b\delta} \text{ degrees}, \\ \theta_2 &= 0.618\theta_1, \\ \psi &= 72^\circ. \end{aligned} \quad (109)$$

It should perhaps be pointed out that the precision of θ_1 , θ_2 , and ψ individually is not critical, since there is no necessity for the coupling to TE_{21} and TE_{31} to be exactly zero, so long as it is reasonably small. One should, however, strive to make the TE_{01} - TM_{11} coupling as nearly zero as possible, and it is therefore important to satisfy (107) with the greatest possible precision.

Expressions for the dielectric coupling coefficients due to the three-sector compensator may be obtained from those for the single-sector compensator if we merely replace $\sin n\theta/2$, wherever it occurs, by $\sin n\theta_1/2 + 2 \cos n\psi \sin n\theta_2/2$, where n is the angular mode index, as usual. If θ_1 , θ_2 , and ψ satisfy (108), then the coupling to TE_{2m} , TE_{3m} , TM_{2m} , and TM_{3m} is zero, and the mode having the largest value of $2(c+d)/(h_0 - h_1)$ is TE_{12} . As noted earlier, since the fields of TM_{11}

TABLE III — COUPLING COEFFICIENT AND POWER TRANSFER TO TE₁₂ MODE DUE TO THREE-SECTOR COMPENSATOR

Mode	$\beta a = 12.930$			$\beta a = 29.554$		
	c	d	$2 \frac{(c + d)}{(h_0 - h_1)}$	c	d	$2 \frac{(c + d)}{(h_0 - h_1)}$
TE ₁₂	3.7591/ b	-3.0334/ b	2.549 a/b	9.0919/ b	-6.5491/ b	21.59 a/b

and TE₁₂ have the same angular dependence, the coupling to TE₁₂ is independent of the number and arrangement of sectors used in the compensator so, long as the decoupling condition for TM₁₁ is satisfied. Numerical values are given in Table III.

The formula analogous to (106) for the attenuation constant due to dielectric loss in a three-sector compensator is

$$\alpha_d = \frac{\beta^2(1 + \delta) \tan \varphi}{2h_{[01]}} \frac{(\theta_1^\circ + 2\theta_2^\circ)}{360^\circ} \text{ nepers/meter.} \tag{110}$$

As a numerical example, let us design a three-sector compensator for a $\frac{7}{8}$ -inch guide at 5.4 mm. Under the requirement that the TE₀₁ loss due to conversion into TE₁₂ must not be greater than 0.1 db, the minimum bending radius is 7.39 inches. If the angles are *

$$\begin{aligned} \theta_1 &= 60^\circ, \\ \theta_2 &= 30^\circ, \\ \psi &= 75^\circ, \end{aligned}$$

the value of δ should be 0.143, which is not difficult to obtain with foam dielectrics. Assuming a loss tangent of 2×10^{-4} we find that dielectric losses in a 90° bend are about 0.12 db.

As a second example, for a 2-inch guide at 5.4 mm with the same mode conversion criterion, one needs a bending radius of 143.1 inches or 11.92 feet. With $\delta = 0.033$, the compensator angles are

$$\begin{aligned} \theta_1 &= 27.6^\circ, \\ \theta_2 &= 16.4^\circ, \\ \psi &= 72.5^\circ; \end{aligned}$$

* It is not practicable to use larger angles, because if $\theta_1 > 60^\circ$ portions of the outer sectors counteract the effect of the rest of the compensator on TM₁₁, and dielectric losses make the design inefficient.

and if $\tan \varphi = 5 \times 10^{-5}$, the dielectric loss in a 90° bend is about 0.25 db.

2.4 *Can Dissipation Be Used to Discourage Spurious Modes?*

It was shown in Section 2.1 that the effect of markedly increasing the attenuation constant of one of two coupled transmission lines is to reduce the over-all attenuation of a wave introduced on the other line. One might wonder whether it would be practicable to decrease the permissible radius of a compensated bend by introducing loss into the spurious modes. The answer is "No", at least for guides large enough to propagate 200 to 300 modes at the operating wavelength. One simply cannot get the required magnitude of loss into the spurious modes without simultaneously introducing intolerable loss into TE_{01} . A numerical example will make this clear.

We found in Section 2.3.3 that with a three-sector compensator in 2-inch guide at 5.4 mm it would be possible to negotiate a bend of radius about 12 feet with a maximum loss of 0.1 db by mode conversion to TE_{12} (the worst spurious mode). Let us now ask what the attenuation constant of TE_{12} would have to be if we wished to transmit around a bend of radius 6 feet with a three-sector compensator, and have the mode conversion loss suffered by TE_{01} not greater than 0.1 db in a 90° bend. Preparing to substitute into (80) of Section 2.1, we have the following values:

$$b = 72 \text{ inches,}$$

$$\kappa_{[12]} = c_{[12]} + d_{[12]} = 2.54/b = 0.0353 \text{ in}^{-1},$$

$$z = \pi b/2 = 113.1 \text{ inches,}$$

$$\beta_0 - \beta_1 \approx h_0 - h_1 = 0.236 \text{ radians/inch,}$$

$$N = 0.1 \text{ db} = 0.0115 \text{ nepers.}$$

From (80) we get

$$\alpha_{[12]} - \alpha_{[01]} = 12.2 \text{ nepers/inch} \approx 4200 \text{ db/meter.}$$

Since the maximum TE_{12} attenuation which can be achieved in a 2-inch guide by a mode filter which transmits TE_{01} freely is of the order of 10 db/meter,* the value of $\alpha_{[12]}$ called for by the above calculation is obviously out of the question.

* This estimate is based on calculations described in Reference 6 for modes in a helix surrounded by a lossy sheath; but it is doubtful that much greater loss could be produced by other types of filter, such as resistance card "killers".

It should be noted that a moderate amount of loss in the spurious modes may be worse than none, so far as the effect on TE_{01} is concerned. Miller¹⁴ has shown that the total power dissipated in the system goes through a maximum when $(\alpha_1 - \alpha_0)/\kappa \approx 2$. It appears that $\alpha_1 - \alpha_0$ must exceed κ by a couple of orders of magnitude before the loss in the driven line (i.e., TE_{01}) becomes really small, if we are counting on dissipation to counteract the coupling to spurious modes.

Since by use of the compensator we are attempting to make the $TE_{01} - TM_{11}$ coupling coefficient zero, we may expect that this coefficient, if not exactly zero, will at least be small compared to the coupling coefficients of spurious modes such as TE_{12} . Because $\kappa_{(11)}$ is very small, it may be that a practicable amount of loss in the TM_{11} mode would improve the performance of the bend. But in view of the preceding paragraph we must be careful, when introducing loss into TM_{11} , not to introduce the wrong amount of loss into some spurious mode which has a larger coupling coefficient to TE_{01} .

ACKNOWLEDGMENTS

I am indebted to S. E. Miller, A. P. King, and J. A. Young for stimulating discussions and several helpful suggestions relating to this work.

APPENDIX

Compensation of a Gradual Bend by a Dielectric Insert in the Adjacent Straight Pipe

We shall discuss briefly three different ways of transmitting the TE_{01} mode around a plain (i.e., air-filled) bend with the aid of dielectric mode transducers inserted into the straight sections of guide on one or both sides of the bend. The first two methods involve converting the TE_{01} mode to a normal mode of the bend and reconverting to TE_{01} on the other side.⁹ In the third method, the input to the bend is pure TE_{01} , and the output mixture of TE_{01} and TM_{11} , whatever it may be, is reconverted to TE_{01} by a dielectric transducer.

A.1 *The TM_{11}' Normal Mode Solution*

One of the normal modes of the bend is a pure TM_{11} mode (TM_{11}') which is polarized at right angles to the TM_{11} mode (TM_{11}'') that the bend couples to TE_{01} . Clearly if one has a transducer in a straight guide which converts TE_{01} entirely to TM_{11} , it is a mere matter of rotating the

transducer about the guide axis to insure that the polarization which enters the bend is TM_{11}' . We shall design such a transducer using a dielectric sector in a straight pipe.

From Section 2.1, for complete power transfer from TE_{01} to TM_{11} we must have

$$\beta_{[01]} - \beta_{(11)} = 0; \quad (111)$$

the transfer then takes place in a distance

$$l = \pi/2 \mid \kappa_{(11)} \mid. \quad (112)$$

The modified phase constants $\beta_{[01]}$ and $\beta_{(11)}$ for a single dielectric sector of angle θ are given by (103) of Section 2.3.2. Substituting these values into (111), we find that the only condition under which it is satisfied is

$$\begin{aligned} \sin \theta &= 0, \\ \theta &= 180^\circ. \end{aligned} \quad (113)$$

The transducer must therefore be a half cylinder. From (104) we have

$$\kappa_{(11)} = d_{(11)} = -0.12066 \beta \delta, \quad (114)$$

and so (112) gives for the length of the transducer,

$$l = 2.072 \lambda / \delta. \quad (115)$$

The $\text{TE}_{01} - \text{TM}_{11}'$ transducer should be placed on either side of the diametral plane of the straight guide which lies in the plane of the bend. An exactly similar transducer on the other side of the bend will reconvert TM_{11}' into TE_{01} . Since TE_{01} and TM_{11} have the same velocity in a straight guide, the transducer can be made of a number of sections with arbitrary spacing and of total length l ; but in practice one will not wish to have too long a run of TM_{11} in the empty guide because of the higher heat losses of this mode.

A.2 The $\text{TE}_{01} \pm \text{TM}_{11}''$ Normal Mode Solution

If we write the coupled wave equations for TE_{01} and TM_{11}'' in a plain bend in the form

$$\begin{aligned} \frac{da_0}{dz} + iha_0 + ica_1 &= 0, \\ ica_0 + \frac{da_1}{dz} + iha_1 &= 0, \end{aligned} \quad (116)$$

where

$$\begin{aligned} h &= h_{[01]} = h_{(11)}, \\ c &= c_{(11)}, \end{aligned} \quad (117)$$

it is evident that we can add and subtract to get the equivalent pair of equations:

$$\begin{aligned} \frac{d(a_0 + a_1)}{dz} + i(h + c)(a_0 + a_1) &= 0, \\ \frac{d(a_0 - a_1)}{dz} + i(h - c)(a_0 - a_1) &= 0. \end{aligned} \quad (118)$$

Hence the normal modes of the curved guide are the combinations $a_0 \pm a_1$, or $TE_{01} \pm TM''$.

In order to launch only the normal mode $TE_{01} + TM''$ at the input of the bend, the amplitude of the other mode must be zero, or $a_0 = a_1$. Similarly, to launch only the mode $TE_{01} - TM''$, the condition is $a_0 = -a_1$. Hence the output of the normal mode transducer of length l , say, must be

$$a_0(l) = \pm a_1(l). \quad (119)$$

We return to the solution of the coupled line equations in Section 2.1 and write

$$\begin{aligned} \kappa &= d_{(11)} = d, \\ \gamma_0 &= i\beta_{[01]} = i\beta_0, \\ \gamma_1 &= i\beta_{(11)} = i\beta_1, \\ r &= is = i\sqrt{(\beta_0 - \beta_1)^2 + 4d^2}. \end{aligned} \quad (120)$$

Then equations (60) of Section 2.1 become:

$$\begin{aligned} a_0(l) &= \left[\cos \frac{1}{2}sl - i \frac{(\beta_0 - \beta_1)}{s} \sin \frac{1}{2}sl \right] e^{-\frac{1}{2}i(\beta_0 + \beta_1)l}, \\ a_1(l) &= -i \frac{2d}{s} \sin \frac{1}{2}sl e^{-\frac{1}{2}i(\beta_0 + \beta_1)l}. \end{aligned} \quad (121)$$

Substituting into (119) and equating real and imaginary parts gives:

$$l = \pi/s, \quad (122)$$

$$|\beta_0 - \beta_1| = |2d|. \quad (123)$$

In view of (103) and (104) of Section 2.3.2, (123) is equivalent to

$$\cos \theta/2 = \frac{\pm \sqrt{1 - \nu^2}}{0.4640\nu^2 + 0.1955}, \quad (124)$$

and (122) becomes

$$l = \frac{1.465\lambda}{\delta |\sin \theta/2|}, \quad (125)$$

where

$$\nu = \lambda/\lambda_c = 3.8317 \lambda/2\pi a \quad (126)$$

is the cutoff ratio for TE_{01} and TM_{11} waves in a straight, empty guide.

A TE_{01} to $TE_{01} \pm TM_{11}$ mode transducer may thus consist of a dielectric sector, attached to the surface of the straight guide on the side nearest the center of curvature of the bend, if the angle of the sector satisfies (124) and the length satisfies (125). However, the condition (124) can be satisfied by a real angle only if

$$0.8483 \leq \lambda/\lambda_c \leq 1; \quad (127)$$

that is, only if the guide is very near cutoff; and the value of θ which satisfies (124) varies rapidly with λ over the above range. This form of normal mode transducer is therefore too narrow band to be of much practical interest.

A.3 A Broadband Compensator

We shall now show how to design a dielectric compensator, in a straight section of guide, which takes the mixture of TE_{01} and TM_{11} put out by an adjacent bend and reconverts it to pure TE_{01} , independent of frequency.

First let us write the solution of (116) for a plain bend in terms of arbitrary initial values $a_0(0)$ and $a_1(0)$. We have

$$\begin{aligned} a_0(z) + a_1(z) &= [a_0(0) + a_1(0)]e^{-ihz-icz}, \\ a_0(z) - a_1(z) &= [a_0(0) - a_1(0)]e^{-ihz+icz}, \end{aligned} \quad (128)$$

and hence

$$\begin{aligned} a_0(z) &= [a_0(0) \cos cz - ia_1(0) \sin cz]e^{-ihz}, \\ a_1(z) &= [-ia_0(0) \sin cz + a_1(0) \cos cz]e^{-ihz}. \end{aligned} \quad (129)$$

The bend may be compensated with a dielectric-loaded straight guide,

provided that* $\beta_0 = \beta_1$ in the straight guide (this may be arranged, for example, by using a half-cylinder of dielectric), and provided that the length of the compensator and the coupling coefficient d are properly chosen. The amplitudes of the two modes in the straight guide, in terms of arbitrary initial values $a_0(0)$ and $a_1(0)$, are

$$\begin{aligned} a_0(z) &= [a_0(0) \cos dz - ia_1(0) \sin dz]e^{-i\beta_0 z}, \\ a_1(z) &= [-ia_0(0) \sin dz + a_1(0) \cos dz]e^{-i\beta_0 z}. \end{aligned} \quad (130)$$

Now suppose that the length of the bend is l_1 and the length of the compensator l_2 , and take the origin of z at the input to the bend. A pure TE_{01} input is represented by

$$\begin{aligned} a_0(0) &= 1, \\ a_1(0) &= 0, \end{aligned} \quad (131)$$

and so, applying (129) and (130) in succession,

$$\begin{aligned} a_0(l_1) &= \cos cl_1 e^{-ihl_1}, \\ a_1(l_1) &= -i \sin cl_1 e^{-ihl_1}, \end{aligned} \quad (132)$$

$$\begin{aligned} a_0(l_1 + l_2) &= \cos (cl_1 + dl_2)e^{-i(hl_1 + \beta_0 l_2)}, \\ a_1(l_1 + l_2) &= -i \sin (cl_1 + dl_2)e^{-i(hl_1 + \beta_0 l_2)}. \end{aligned} \quad (133)$$

The condition that all the power be in TE_{01} at the output of the compensator at every frequency is

$$cl_1 + dl_2 = 0, \quad (134)$$

or

$$dl_2 = -0.18454 \beta a l_1 / b, \quad (135)$$

on referring to equation (86) of Section 2.2 for the value of c .

A convenient form of compensator would be a half cylinder of dielectric whose diametral plane is perpendicular to the plane of the bend. The coupling coefficient of the half cylinder is given by (104), and the condition for complete compensation becomes

$$l_2 \delta = 1.5295 l_1 a / b. \quad (136)$$

The most easily adjustable parameter is probably the length l_2 of the compensator.

* The necessity for $\beta_0 = \beta_1$ is apparent if we consider that under certain conditions the bend may put out a pure TM_{11} mode, and complete reconversion is possible only if the compensator has $\beta_0 = \beta_1$.

We have analyzed the compensator as if it were all on one side of the bend; but it may evidently be divided into sections of total length l_2 which are distributed arbitrarily on both sides of the bend. An obvious configuration would be to put a section of length $l_2/2$ on each side of the bend immediately adjacent to it.

Limitations on the usefulness of this kind of compensator will be dielectric losses and mode conversions. The former can presumably be reduced as the dissipation of available dielectrics is reduced. Mode conversions could be calculated by the general methods of Part I, but one would have to work out the values of the coupling coefficients between TM_{11} and all other modes, which has not yet been done. In any case it is likely that the minimum tolerable bending radius would be no less than for the within-the-bend compensators discussed earlier in this paper.

REFERENCES

1. M. Jouguet, *Cables & Transmission*, **1**, pp. 133-153, 1947.
2. S. O. Rice, unpublished memorandum.
3. W. J. Albersheim, *B. S. T. J.*, **28**, pp. 1-32, January, 1949.
4. S. E. Miller, *Proc. I. R. E.*, **40**, pp. 1104-1113, September, 1952.
5. H. G. Unger, pp. 1253-1278, this issue.
6. S. A. Morgan and J. A. Young, *B. S. T. J.*, **35**, pp. 1347-1384, November, 1956.
7. S. A. Schelkunoff, *B. S. T. J.*, **31**, pp. 784-801, July, 1952.
8. S. E. Miller, *B. S. T. J.*, **33**, pp. 661-719, May, 1954; especially pp. 677-692.
9. Reference 4, p. 1110.
10. S. A. Schelkunoff, *Electromagnetic Waves*, D. Van Nostrand Co., Inc., New York, 1943, pp. 12, 94.
11. Reference 7, pp. 786-791.
12. Reference 7, pp. 787-788.
13. S. O. Rice, *B. S. T. J.*, **27**, pp. 305-349, April, 1948; especially p. 348.
14. Reference 8, pp. 692-693, Figs. 31 and 32.

Circular Electric Wave Transmission in a Dielectric-Coated Waveguide

By H. G. UNGER

(Manuscript received January 9, 1957)

The mode conversion from the TE_{01} wave to the TM_{11} wave which normally occurs in a curved round waveguide may be reduced by applying a uniform dielectric coat a few mils thick to the inner wall of the waveguide. Such a dielectric coat changes the phase constant of the TM_{11} wave without affecting appreciably the phase and attenuation constants of the TE_{01} wave. Thus, relatively sharp bends can be negotiated to change the direction of the line or deviations from straightness can be tolerated. For each of these two cases a different optimum thickness of the dielectric coat keeps the TE_{01} loss at a minimum. At 5.4-mm wavelength a bending radius of 8 ft in a $\frac{7}{8}$ -inch pipe or 50 ft in a 2-inch pipe can be introduced with 0.2 db mode conversion loss. Deviations from straightness corresponding to an average bending radius of 300 ft can be tolerated in a 2-inch pipe at 5.4-mm wavelength with an increase in TE_{01} attenuation of 5 per cent. Serpentine bends caused by equally spaced supports of the pipe may, however, increase the mode conversion and cause appreciable loss at certain discrete frequencies. Compared with the plain waveguide the dielectric-coated guide behaves more critically in such serpentine bends. The mode conversion from TE_{01} to the TE_{02} , TE_{03} , $\dots$ waves at transitions from a plain waveguide to the dielectric-coated guide is usually very small.

I. INTRODUCTION

In a curved section of cylindrical waveguide the circular electric wave couples to the TE_{11} , TE_{12} , TE_{13} $\dots$ modes and to the TM_{11} mode.¹ The coupling to the TM_{11} mode presents the most serious problem since the TE_{01} and TM_{11} modes are degenerate, in that they have equal phase velocities in a perfectly conducting straight guide. In a bend all TE_{01} power introduced at the beginning will be converted to the TM_{11} power at odd multiples of a certain critical bending angle. One can reduce this complete power transfer by removing the degeneracy between TE_{01} and

TM₁₁ modes. The finite conductivity of the walls introduces a slight difference in the propagation constants and in a 2-inch pipe at 5.4 mm wavelength a bending radius of a few miles can be tolerated with about double the attenuation constant of the TE₀₁ wave. To get more difference in the phase constants of the TM₁₁ and TE₀₁ modes, one might consider a dielectric layer next to the wall of the waveguide. Since the electric field intensity of the TE₀₁ mode goes to zero at the wall but the electric field intensity of the TM₁₁ mode has a large value there, one might expect a larger effect of the dielectric layer on the propagation characteristics of the TM₁₁ wave than on the TE₀₁ wave.

In doing this, however, one has to be aware of the influence the dielectric layer will have on the propagation characteristics of the TE_{1_m} modes which also couple to the TE₀₁ wave in curved sections. The TE₁₂ wave couples most strongly to the TE₀₁ wave and of the TE_{1_m} family it is the next above the TE₀₁ wave in phase velocity. With the dielectric layer one has to expect this difference in phase velocity to be decreased and consequently the mode conversion to the TE₁₂ wave to be increased.

In the next section we will solve the characteristic equation of the cylindrical waveguide with a dielectric layer for the normal modes and arrive at approximate formulas for the phase constants. We will also compute the increase in attenuation of the normal modes as caused by the dielectric losses in the layer. The change in wall current losses as caused by the dielectric layer is of importance here only for the TE₀₁ wave and we will calculate it only for this wave.

In order to know what bending radii can be tolerated with a dielectric coat, we have to evaluate the coupled wave theory² for small differences in propagation constants between the coupled waves. This will be done in Section III. Especially we will consider serpentine bends which occur in any practical line with discrete supports. In that situation, at certain critical frequencies, when the supporting distance is a multiple of the beat wavelength between TE₀₁ and a particular coupled wave, we will have to expect serious effects on the propagation constant of the TE₀₁ wave.

In Section IV we will combine the results of Sections II and III and establish formulas and curves for designing circular electric waveguides with a dielectric coat. We will distinguish there between two different applications. The first case is to negotiate uniform bends of as small a radius as possible and the second is to tolerate large bending radii which may occur in a normally straight line.

At transitions from plain waveguide to the dielectric-coated guide, power of an incident TE₀₁ wave will be partly converted into higher

TE_{om} waves. In Section V we shall calculate the power level in the TE_{om} waves resulting from this conversion.

II. THE NORMAL MODES OF THE DIELECTRIC-COATED WAVEGUIDE

The waveguide structure under consideration is shown in Fig. 1. To find the various normal modes existing in this structure, Maxwell's equations have to be solved in circular cylindrical coordinates in the air-filled region 1 and dielectric-filled region 2. The boundary conditions are: equal tangential components of the electric and magnetic field intensities at the boundary ($r = b$) between regions 1 and 2 and, assuming infinite conductivity of the walls, zero tangential component of the electric field at the walls ($r = a$). Upon introducing the general solutions into these boundary conditions, we get a homogeneous system of four linear equations in the amplitude factors. Non-trivial solutions of this system require the coefficient determinant to be zero. This condition is called the characteristic equation. Solutions of the characteristic equation represent the propagation constants of the various modes. These calculations have been carried out elsewhere,^{3, 4} and the characteristic equation arrived at there has the following form:

$$n^2 \left[\frac{1}{x_1^2} - \frac{1}{x_2^2} \right]^2 - \rho^2 \frac{x_2^2 - x_1^2}{x_2^2 - \epsilon x_1^2} \left[\frac{1}{x_1} \frac{J_n'(\rho x_1)}{J_n(\rho x_1)} + \frac{\epsilon}{\rho x_2^2} \frac{W_n(x_2, \rho x_2)}{U_n(x_2, \rho x_2)} \right] \cdot \left[\frac{1}{x_1} \frac{J_n'(\rho x_1)}{J_n(\rho x_1)} + \frac{1}{\rho x_2^2} \frac{V_n(x_2, \rho x_2)}{Z_n(x_2, \rho x_2)} \right] = 0. \quad (1)$$

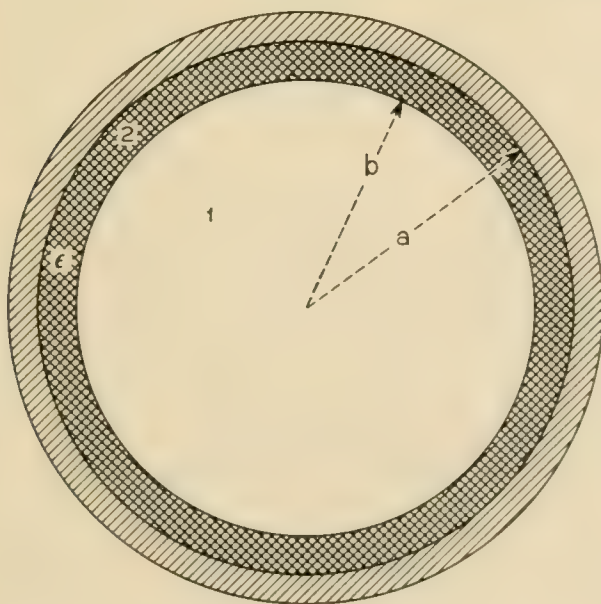


FIG. 1 — The dielectric-coated waveguide; $\rho = b/a$, $\delta = 1 - \rho = (a - b)/a$.

Here we have used the following definitions:

$$\begin{aligned} U_n(x, \rho x) &= J_n(\rho x)N_n(x) - N_n(\rho x)J_n(x), \\ V_n(x, \rho x) &= \rho x^2[J_n'(\rho x)N_n'(x) - N_n'(\rho x)J_n'(x)], \\ W_n(x, \rho x) &= \rho x[J_n(x)N_n'(\rho x) - J_n'(\rho x)N_n(x)], \\ Z_n(x, \rho x) &= x[J_n'(x)N_n(\rho x) - J_n(\rho x)N_n'(x)]. \end{aligned} \quad (2)$$

J_n and N_n are cylinder functions of the first and second kind, respectively, and the prime marks differentiation with respect to the argument. The other symbols are:

$$\begin{aligned} \rho &= b/a \text{ ratio of radii,} \\ x_1 &= \xi_1 a, \\ x_2 &= \xi_2 a. \end{aligned} \quad (3)$$

The relative dielectric constant of region 2 is indicated by ϵ which is assumed to be complex to take dielectric losses into account. The radial propagation constants ξ_1 and ξ_2 are related to the axial propagation constant γ and the free-space propagation constant $k = 2\pi/\lambda$ by

$$\begin{aligned} \xi_1^2 &= k^2 + \gamma^2, \\ \xi_2^2 &= \epsilon k^2 + \gamma^2. \end{aligned} \quad (4)$$

The circumferential order of the solution is indicated by n .

For $n = 0$, equation (1) splits into the two equations

$$\frac{1}{x_1} \frac{J_0'(\rho x_1)}{J_0(\rho x_1)} + \frac{\epsilon}{\rho x_2^2} \frac{W_0(x_2, \rho x_2)}{U_0(x_2, \rho x_2)} = 0, \quad (5)$$

and

$$\frac{1}{x_1} \frac{J_0'(\rho x_1)}{J_0(\rho x_1)} + \frac{1}{\rho x_2^2} \frac{V_0(x_2, \rho x_2)}{Z_0(x_2, \rho x_2)} = 0 \quad (6)$$

representing the TM_{0m} and TE_{0m} waves, respectively.

Except for $n = 0$ the solutions of (1) and the modes in the waveguide do not have transverse character as in the case of a uniformly filled waveguide. They are hybrid modes. However, it is reasonable to label them as TE_{nm} or TM_{nm} , according to the limit which they approach as the dielectric layer becomes very thin.⁵

Modes in the dielectric-coated guide with a very thin coat may be treated as perturbed TE_{nm} or TM_{nm} modes of the ideal circular waveguide. The perturbation terms are found by expanding (1) for small

values of $\delta = 1 - \rho$. This is done in Appendix 1. The perturbation of the propagation constants for the various normal modes is:

$$\text{TM}_{nm} \frac{\Delta\gamma}{\gamma_{nm}} = \frac{\epsilon - 1}{\epsilon} \delta, \quad (7)$$

$$\text{TE}_{nm} \frac{\Delta\gamma}{\gamma_{nm}} = \frac{n^2}{p_{nm}^2 - n^2} \frac{\epsilon - 1}{\epsilon(1 - \nu_{nm}^2)} \delta, \quad (8)$$

$$\text{TE}_{om} \frac{\Delta\gamma}{\gamma_{om}} = \frac{\epsilon - 1}{1 - \nu_{om}^2} \frac{p_{om}^2}{3} \delta^3. \quad (9)$$

In these equations p_{nm} is the m th root of $J_n(x) = 0$ for TM_{nm} waves and the m th root of $J'_n(x) = 0$ for TE_{nm} waves. Furthermore, $\nu_{nm} = \lambda/\lambda_{cnm}$ where $\lambda_{cnm} = 2\pi a/p_{nm}$.

We note that the change in propagation constant is of first order in δ for the TM_{nm} and TE_{nm} waves with $n \neq 0$, but of third order in δ for TE_{om} waves.

With $\epsilon = \epsilon' - j\epsilon''$ the perturbation of the propagation constant splits into phase perturbation $\Delta\beta$ and dielectric attenuation α_D . For a low loss dielectric we may assume $\epsilon'' \ll \epsilon'$ and get:

$$\begin{aligned} \text{TM}_{nm} \frac{\Delta\beta}{\beta_{nm}} &= \frac{\epsilon' - 1}{\epsilon'} \delta, \\ \text{TE}_{nm} \frac{\Delta\beta}{\beta_{nm}} &= \frac{n^2}{p_{nm}^2 - n^2} \frac{\epsilon' - 1}{\epsilon'(1 - \nu_{nm}^2)} \delta, \\ \text{TE}_{om} \frac{\Delta\beta}{\beta_{om}} &= \frac{p_{om}^2}{3} \frac{\epsilon' - 1}{1 - \nu_{om}^2} \delta^3, \\ \text{TM}_{nm} \frac{\alpha_D}{\beta_{nm}} &= \frac{\epsilon''}{\epsilon'^2} \delta, \\ \text{TE}_{nm} \frac{\alpha_D}{\beta_{nm}} &= \frac{n^2}{p_{nm}^2 - n^2} \frac{\epsilon''}{\epsilon'^2(1 - \nu_{nm}^2)} \delta, \\ \text{TE}_{om} \frac{\alpha_D}{\beta_{nm}} &= \frac{p_{om}^2}{3} \frac{\epsilon''}{1 - \nu_{om}^2} \delta^3. \end{aligned} \quad (10)$$

Unfortunately the range in which (7), $\dots$ (10) are good approximations is rather limited. Actually

$$\frac{1 - \nu_{nm}^2}{\nu_{nm}} \cdot \frac{2\pi a}{\lambda} \cdot \frac{\Delta\beta}{\beta_{nm}}$$

has to be small compared to unity, at least not larger than 0.1. In the

case of the TM_{11} wave in a 2-inch pipe at $\lambda = 5.4$ mm this limit is reached with $\Delta\beta/\beta_{11} = 0.45 \times 10^{-3}$ or $\delta = 0.75 \times 10^{-3}$ with $\epsilon = 2.5$.

To evaluate (1) beyond this limit we may use the approximations:

$$\begin{aligned} \frac{1}{\rho x} \frac{W_n(x, \rho x)}{U_n(x, \rho x)} &= \cot (1 - \rho)x = \cot \delta x, \\ \frac{1}{\rho x} \frac{V_n(x, \rho x)}{Z_n(x, \rho x)} &= -\tan (1 - \rho)x = -\tan \delta x. \end{aligned} \quad (11)$$

The approximations (11) require $x \gg n$, and $\delta x < \pi/2$, as shown in Appendix I. These requirements are usually satisfied for the lower order modes in a multimode waveguide with a thin dielectric layer.

Equation (1) has been evaluated for the TM_{11} mode using the approximations (11) and a method which is described in Appendix I. The results are shown in Fig. 2 for $\epsilon = 2.5$ and several values of a/λ .

In introducing a dielectric coat we have to be aware of the change in attenuation constant the TE_{01} mode will suffer. Not only the loss factor of the dielectric material will cause additional losses, but the concentration of more field energy into the dielectric will increase the wall currents and so the wall current losses.

To calculate the wall current attenuation of the TE_{01} mode in the round waveguide with the dielectric layer, we proceed in the usual

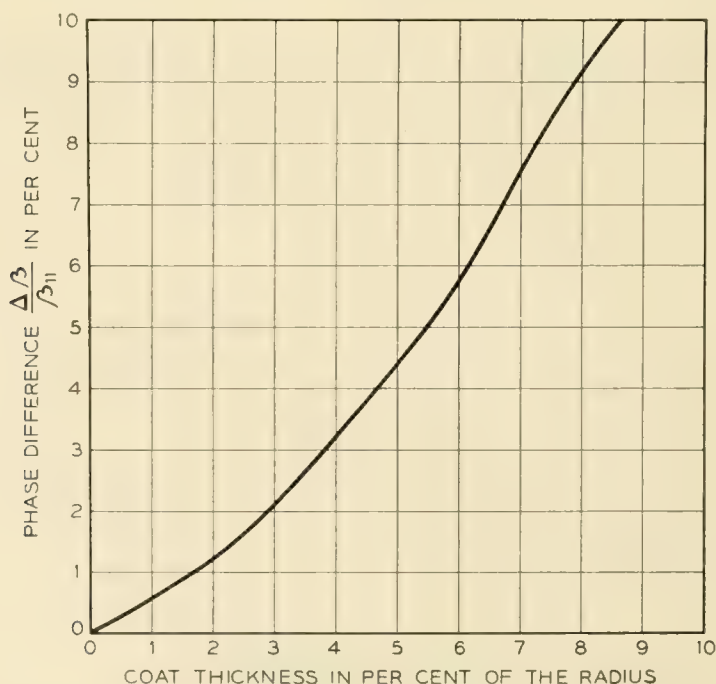


Fig. 2(a) — Change in phase constant of the TM_{11} wave in the dielectric-coated waveguide. $\epsilon = 2.5$; $a/\lambda = 1.03$.

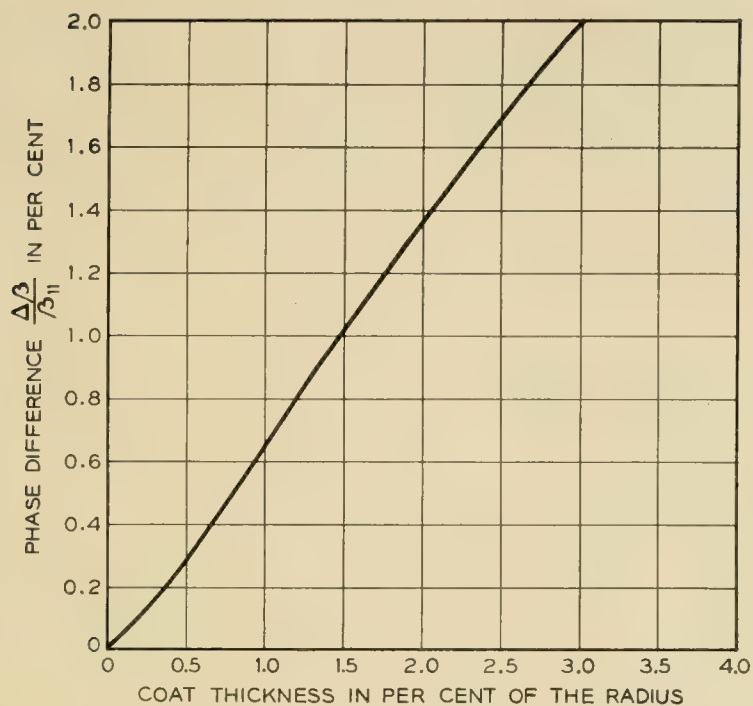


Fig. 2(b) — Change in phase constant of the TM_{11} wave in the dielectric-coated waveguide. $\epsilon = 2.5$; $a/\lambda = 2.06$.

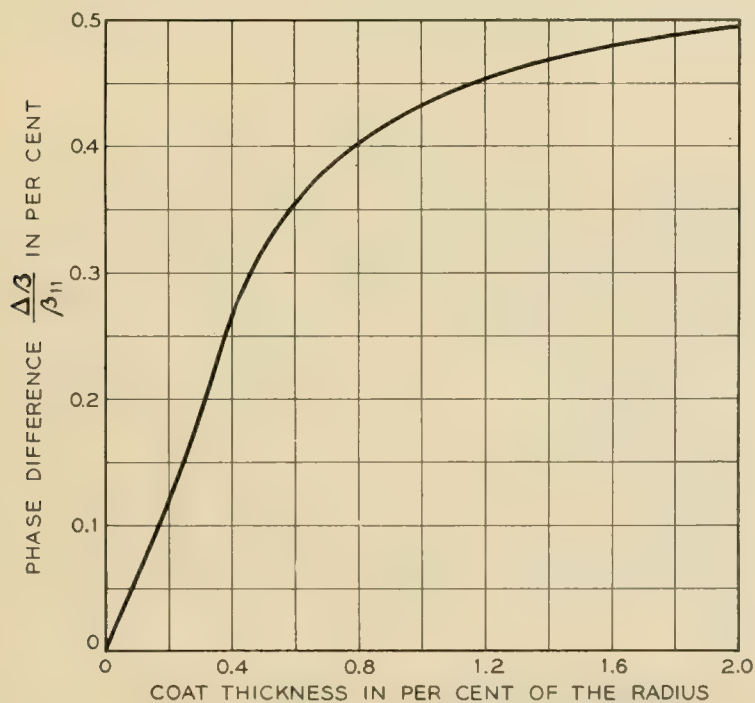


Fig. 2(c) — Change in phase constant of the TM_{11} wave in the dielectric-coated waveguide. $\epsilon = 2.5$; $a/\lambda = 4.70$.

manner. Assuming the conductivity to be high though finite and the loss factor of the dielectric material to be small, we take the field pattern and wall currents of the lossless case and calculate the total transmitted power P and the power P_M absorbed per unit length of the waveguide by the metal walls of finite conductivity. The wall current attenuation α_M is then given by

$$\alpha_M = \frac{1}{2} \frac{P_M}{P}. \quad (12)$$

The result of this calculation as carried through in Appendix II is:

$$\frac{\alpha_M}{\alpha_{om}} = \frac{x_2^2}{p_{om}^2} \sqrt{\frac{x_2^2 - x_1^2}{x_2^2 - \epsilon x_1^2}} \left\{ 1 + \frac{\pi^2}{2} \rho x_2 \left[\frac{x_2^2}{x_1^2} - 1 \right] \right. \\ \left. \cdot \left[U_1(x_2, \rho x_2) - \frac{\rho x_2}{2} R_1(x_2, \rho x_2) \right] R_1(x_2, \rho x_2) \right\}^{-1}. \quad (13)$$

In this expression α_M is related to the TE_{om} attenuation constant α_{om} of the plain waveguide.

For $1 - \rho = \delta \ll 1$ we introduce the series expansions of the functions $U_1(x_2, \rho x_2)$ and $R_1(x_2, \rho x_2)$,

$$\frac{\Delta \alpha_M}{\alpha_{om}} = (\epsilon - 1) \frac{p_{om}^2}{\nu_{om}^2} \delta^2. \quad (14)$$

Here $\Delta \alpha_M$ is defined as the change in wall current attenuation compared to the attenuation in the plain waveguide.

III. PROPERTIES OF COUPLED TRANSMISSION LINES

Wave propagation in gentle bends of a round waveguide can be described in terms of normal modes of the straight guide.¹ The bend causes coupling between the normal modes. The TE_{01} wave couples to the TM_{11} wave and to the TE_{1n} waves and the propagation in the bend is described by an infinite set of simultaneous linear differential equations. An adequate approximate treatment is to consider only coupling between TE_{01} and one of the spurious modes at a time. Furthermore, only the forward waves need to be considered, since the relative power coupled from the forward waves into the backward waves is quite small. Thus, the infinite set of equations reduces to the well known coupled line equations:²

$$\begin{aligned} \frac{dE_1}{dz} + \gamma_1 E_1 - jcE_2 &= 0, \\ \frac{dE_2}{dz} + \gamma_2 E_2 - jcE_1 &= 0, \end{aligned} \quad (15)$$

in which

$E_{1,2}(z)$ = wave amplitudes in mode 1 (here always TE_{01}) and mode 2 (TM_{11} or one of the TE_{1n}), respectively;

$\gamma_{1,2}$ = propagation constant of mode 1 and 2, respectively (the small perturbation of $\gamma_{1,2}$ caused by the coupling may be neglected here); and

c = coupling coefficient between modes 1 and 2.

Subject to the initial conditions:

$$E_1(0) = 1, \quad E_2(0) = 0,$$

the solution of (15) is:

$$\begin{aligned} E_1 &= \frac{1}{2} \left[1 + \frac{\Delta\gamma}{\sqrt{\Delta\gamma^2 - 4c^2}} \right] e^{-\Gamma_1 z} + \frac{1}{2} \left[1 - \frac{\Delta\gamma}{\sqrt{\Delta\gamma^2 - 4c^2}} \right] e^{-\Gamma_2 z}, \\ E_2 &= \frac{j c}{\sqrt{\Delta\gamma^2 - 4c^2}} [e^{-\Gamma_2 z} - e^{-\Gamma_1 z}], \end{aligned} \quad (16)$$

where $\Delta\gamma = \gamma_1 - \gamma_2$ and $\Gamma_{1,2} = \frac{1}{2}[\gamma_1 + \gamma_2 \pm \sqrt{\Delta\gamma^2 - 4c^2}]$. Γ_1 and Γ_2 are the propagation constants of the two coupled line normal modes. Both coupled line normal modes are excited by the initial conditions. For $|c/\Delta\gamma| \ll 1$, equations (16) can be simplified:

$$\begin{aligned} E_1 &= \left[1 - 2 \frac{c^2}{\Delta\gamma^2} \sinh \frac{1}{2} \Delta\gamma z e^{\frac{1}{2} \Delta\gamma z} \right] e^{-(\gamma_1 - (c^2/\Delta\gamma))z}, \\ E_2 &= j \frac{2c}{\Delta\gamma} \sinh \frac{1}{2} \Delta\gamma z e^{-\frac{1}{2}(\gamma_1 + \gamma_2)z}. \end{aligned} \quad (17)$$

We are concerned with a difference in phase constant which is much larger than the difference in attenuation constant. Consequently in $\Delta\gamma = j\Delta\beta + \Delta\alpha$ we have $|\Delta\beta| \gg |\Delta\alpha|$ and we may write:

$$\begin{aligned} E_1 &= \left[1 + j \frac{2c^2}{\Delta\beta^2} \sin \frac{1}{2} \Delta\beta z e^{j\frac{1}{2} \Delta\beta z} \right] \\ &\quad \cdot \exp \left[-j \left(\beta_1 + \frac{c^2}{\Delta\beta^2} \Delta\beta \right) z - \left(\alpha_1 - \frac{c^2}{\Delta\beta^2} \Delta\alpha \right) z \right], \\ E_2 &= j \frac{2c}{\Delta\beta} \sin \frac{1}{2} \Delta\beta z \exp \left[-j \frac{1}{2} (\beta_1 + \beta_2) z - \frac{1}{2} (\alpha_1 + \alpha_2) z \right]. \end{aligned} \quad (18)$$

We note that the amplitude E_1 , apart from suffering an attenuation, varies in an oscillatory manner, the maximum being $E_1 = 1$ and the minimum $E_1 = 1 - 2(c^2/\Delta\beta^2)$. Accordingly, the maximum mode conversion loss is given by:

$$20 \log_{10} \frac{E_{1\max}}{E_{1\min}} = 17.35 \frac{c^2}{\Delta\beta^2}. \quad (19)$$

The attenuation constant of E_1 is modified by the presence of the coupled wave. Compared to the uncoupled attenuation constant it has been changed by

$$\frac{\Delta\alpha_c}{\alpha_1} = \frac{c^2}{\Delta\beta^2} \left(\frac{\alpha_2}{\alpha_1} - 1 \right). \quad (20)$$

The amplitude E_2 varies sinusoidally. From our point of view it is an unwanted mode. The power level compared to the E_1 power is

$$20 \log_{10} \frac{E_{2\max}}{E_1} = 20 \log_{10} \frac{2c}{\Delta\beta}. \quad (21)$$

So far we have considered only a constant value of the coupling coefficient, c , corresponding to a uniform bend. The attenuation in such a uniform bend is increased according to (20), and the worst condition we can encounter at the end of the bend is a mode conversion loss (19) and a spurious mode level (21).

A practical case of changing curvature and consequently changing coupling coefficient is the serpentine bend. A waveguide with equally spaced supports deforms into serpentine bends under its own weight. The curve between two particular supports is well known from the theory of elasticity. An analysis of circular electric wave transmission through serpentine bends⁶ shows that mode conversion becomes seriously high at certain critical frequencies when the supporting distance is a multiple of the beat wavelength between the TE_{01} and a particular coupled mode. The beat wavelength is here defined as

$$\lambda_b = \frac{2\pi}{\Delta\beta}. \quad (22)$$

In serpentine bends formed by elastic curves, mode conversion at the critical frequencies causes an increase in TE_{01} attenuation⁶

$$\frac{\Delta\alpha_s}{\alpha_{01}} = - \left[\frac{w}{EI} \frac{c_0}{\Delta\beta^2\alpha_{01}} \right]^2 \frac{\alpha_{01}}{\Delta\alpha}, \quad (23)$$

and a spurious mode level in the particular coupled mode⁶

$$\left| \frac{E_2}{E_1} \right| = \frac{w}{EI} \frac{c_0}{\Delta\beta^2\alpha_{01}} \left| \frac{\alpha_{01}}{\Delta\alpha} \right|, \quad (24)$$

where w = weight per unit length of the pipe,
 E = modulus of elasticity,
 I = moment of inertia,

and c_0 is the factor in the bend coupling coefficient $c = c_0/R$ determined by waveguide dimensions, frequency and the particular mode. R is the bend radius.

Equations (23) and (24) hold only as long as

$$4\Delta\alpha_s \ll |\Delta\alpha| \quad (25)$$

is satisfied. The rate of conversion loss has to be small compared to the difference between the rates of decay for the unwanted mode and TE_{01} amplitudes. If (25) is not satisfied a cyclical power transfer between TE_{01} and the particular coupled mode occurs, and the TE_{01} transmission is seriously distorted.

Another case of changing curvature of the waveguide is random deviation from straightness, which must be tolerated in any practical line. Such deviations from straightness change the curvature only very gradually, and since there is no coupled wave in the dielectric coated guide, which has the same phase constant as the TE_{01} wave, the curvature may be assumed to vary only slowly compared to the difference in phase constants. Under this condition⁷ the normal mode of the straight waveguide, which here is the TE_{01} mode, will be transformed along the gradually changing curvature into the normal mode of the curved waveguide, which is a certain combination of TE_{01} and the coupled modes. This normal mode will always be maintained and no spurious modes will be excited.

The change in transmission loss is therefore given alone by the difference between the attenuation constants of the TE_{01} wave and the normal mode of the curved waveguide. The propagation constants of the normal modes of the coupled lines are Γ_1 and Γ_2 as given by (16). Only the mode with Γ_1 will be excited here and consequently the attenuation difference is given by (20).

IV. CIRCULAR ELECTRIC WAVE TRANSMISSION THROUGH CURVED SECTIONS OF THE DIELECTRIC-COATED GUIDE

In curved sections of the plain waveguide the wave solution has been described in terms of the normal modes of the straight waveguide. In this presentation it has been found that the TE_{01} wave couples to the TM_{11} and TE_{1n} waves in gentle bends. Likewise, the wave solution in curved sections of the dielectric-coated guide can be described in terms of the normal waves of the straight dielectric-coated guide. In a waveguide with a thin dielectric coat the normal waves may be considered as perturbed normal modes of the plain waveguide. Consequently the bend solution of the plain waveguide may be taken as the first order approximation

for the bend solution of the dielectric-coated guide. In this first order approximation the TE_{01} wave of the dielectric-coated guide couples to the TM_{11} and TE_{1n} waves of the dielectric-coated guide and the coupling coefficients are the same as in the bend-solution of the plain waveguide:¹

$$TE_{01} \rightleftharpoons TM_{11} \quad c = 0.18454 \frac{\beta a}{R},$$

$$TE_{01} \rightleftharpoons TE_{11}$$

$$c = \left[\frac{0.09319(\beta a)^2 - 0.84204}{\sqrt{\beta_{01}a\beta_{11}a}} + 0.09319 \sqrt{\beta_{01}a\beta_{11}a} \right] \frac{1}{R},$$

$$TE_{01} \rightleftharpoons TE_{12} \quad (26)$$

$$c = \left[\frac{0.15575(\beta a)^2 - 3.35688}{\sqrt{\beta_{01}a\beta_{12}a}} + 0.15575 \sqrt{\beta_{01}a\beta_{12}a} \right] \frac{1}{R},$$

$$TE_{01} \rightleftharpoons TE_{13}$$

$$c = \left[\frac{0.01376(\beta a)^2 - 0.60216}{\sqrt{\beta_{01}a\beta_{13}a}} + 0.01376 \sqrt{\beta_{01}a\beta_{13}a} \right] \frac{1}{R},$$

β = free-space phase constant.

With the coupling coefficients (26) and the propagation constants as given by (10) and (14) the coupled line equations can be used to compute the TE_{01} transmission through curved sections. Since the absolute value of $c/\Delta\gamma$ is usually small compared to unity, the mode conversion loss is given by (19), the increase in TE_{01} attenuation by (20), and the spurious mode level by (21).

The curves in Figs. 3, 4 and 5 have been calculated for the 2-inch pipe at a wavelength of 5.4 mm. They take into account the effects of the three most seriously coupled modes, TM_{11} , TE_{11} , and TE_{12} .

In very gentle bends of a guide with a very thin coat, coupling effects to the TE_{1m} modes are small and only TM_{11} coupling influences the TE_{01} transmission. The increase in TE_{01} attenuation in such bends is obtained by substituting the expression (10) for the TM_{11} phase difference in (20).

$$\frac{\Delta\alpha_c}{\alpha_{01}} = \frac{0.034}{\nu_{01}^2} \frac{\epsilon'^2}{(\epsilon' - 1)^2} \frac{1}{\delta^2} \frac{a^2}{R^2}. \quad (27)$$

Note that (27) requires $|c/\Delta\gamma| \ll 1$ and consequently $(1/\delta)(a/R) \ll 1$. In Fig. 4, as well as in (27), losses in the dielectric coat have been neglected. Equations (10) show that the dielectric losses are small compared to the wall current losses for a low loss dielectric coat.

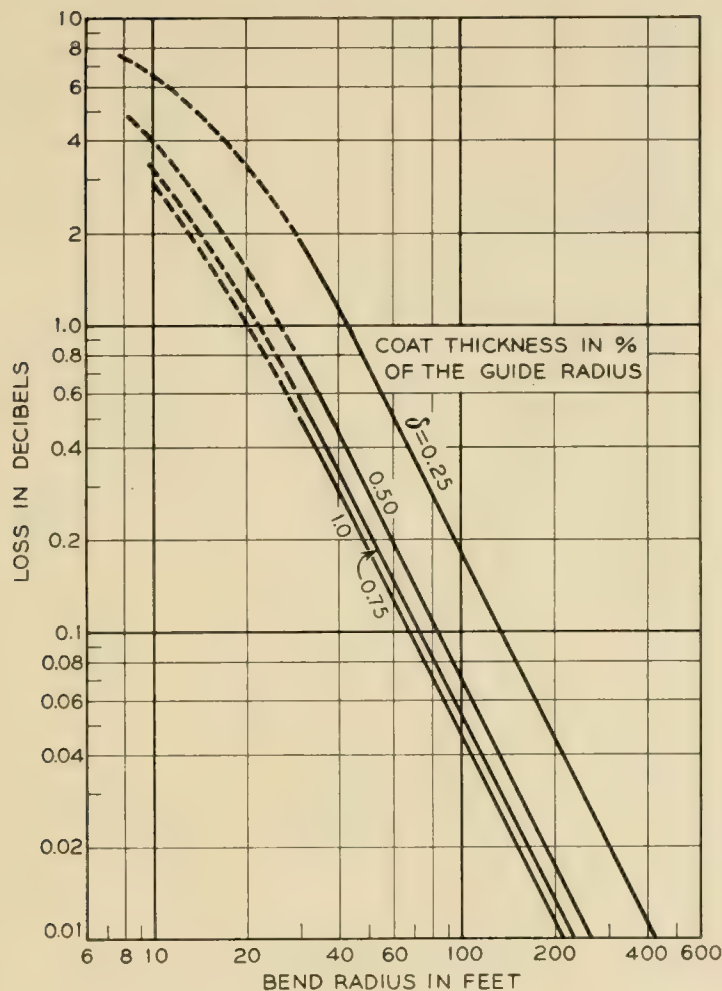


Fig. 3 — Maximum conversion loss of the TE_{01} wave in a uniform bend of the dielectric-coated waveguide. $2a = 2$ inches; $\lambda = 5.4$ mm; $\epsilon = 2.5$.

There are two different applications of the dielectric-coated guide, which require different designs:

1. Intentional bends

These are relatively short sections and the increase in TE_{01} attenuation is usually small compared to the bend loss. Also, a high spurious mode level can be tolerated, because with a mode filter at the end of the bend we can always control the spurious mode level. Consequently, the only limit set for this type of bend is the mode conversion loss of the TE_{01} wave. There is conversion loss mainly to TE_{11} , TE_{12} and TM_{11} . Increasing the phase difference of the TM_{11} wave by making the dielectric coat thicker decreases the phase-difference between TE_{01} and TE_{12} and increases the phase-difference between TE_{01} and TE_{11} . Apparently we get

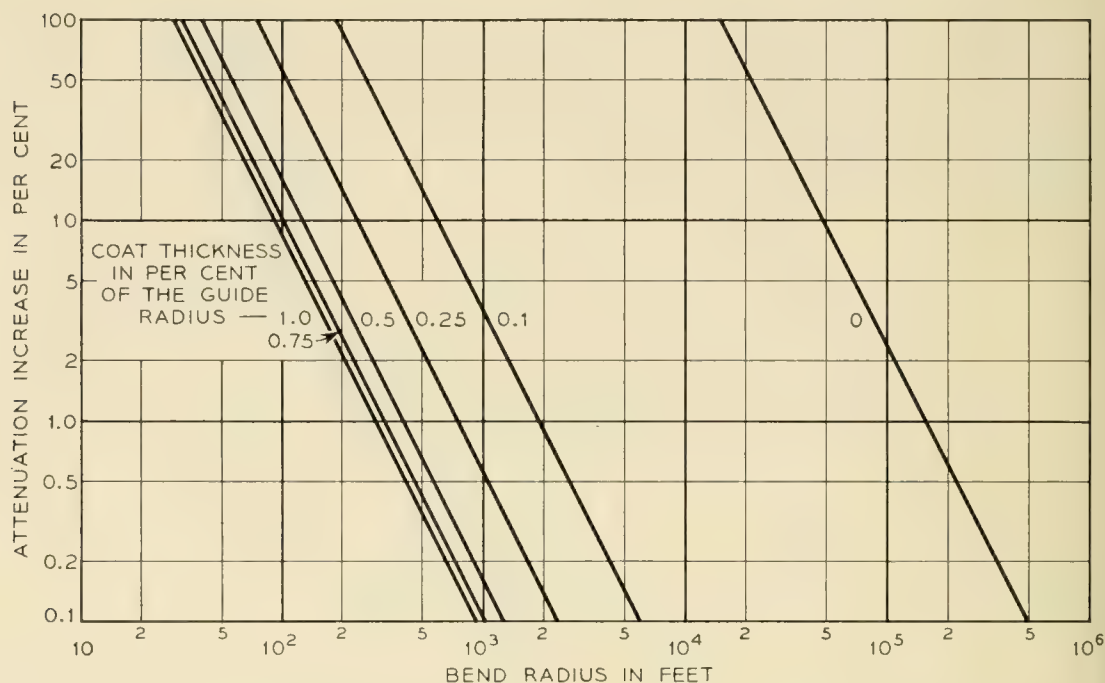


Fig. 4 — Increase in TE_{01} attenuation of a dielectric-coated guide in a uniform bend. $2a = 2$ inches; $\epsilon = 2.5$; $\lambda = 5.4$ mm.

very near to an optimum design with a dielectric coat for which:

$$\left(\frac{\Delta\beta}{c}\right)_{TM_{11}} = \left(\frac{\Delta\beta}{c}\right)_{TE_{12}}. \quad (28)$$

For this condition conversion losses to TM_{11} and TE_{12} are equal, while the conversion loss to TE_{11} is small.

To find values which satisfy (28) we will generally have to solve (1) because the TM_{11} phase difference required by (28) is too large to be calculated with the first order approximation. At a wavelength of $\lambda = 5.4$ mm and a dielectric constant $\epsilon = 2.5$ the optimum thickness of the coat according to (28) is $\delta = 1.25$ per cent in the 2-inch pipe. In Fig. 6 the mode conversion loss in a dielectric-coated guide of this design is plotted versus bending radius.

2. Random deviations from straightness

As mentioned before, random deviations from straightness change the curvature only gradually and only one normal mode propagates. Mode conversion loss and spurious mode level are very low. The normal mode attenuation depends on the curvature. The increase in normal mode attenuation as caused by curvature is obtained by adding the attenuation terms (20) of the various straight guide modes which are contained

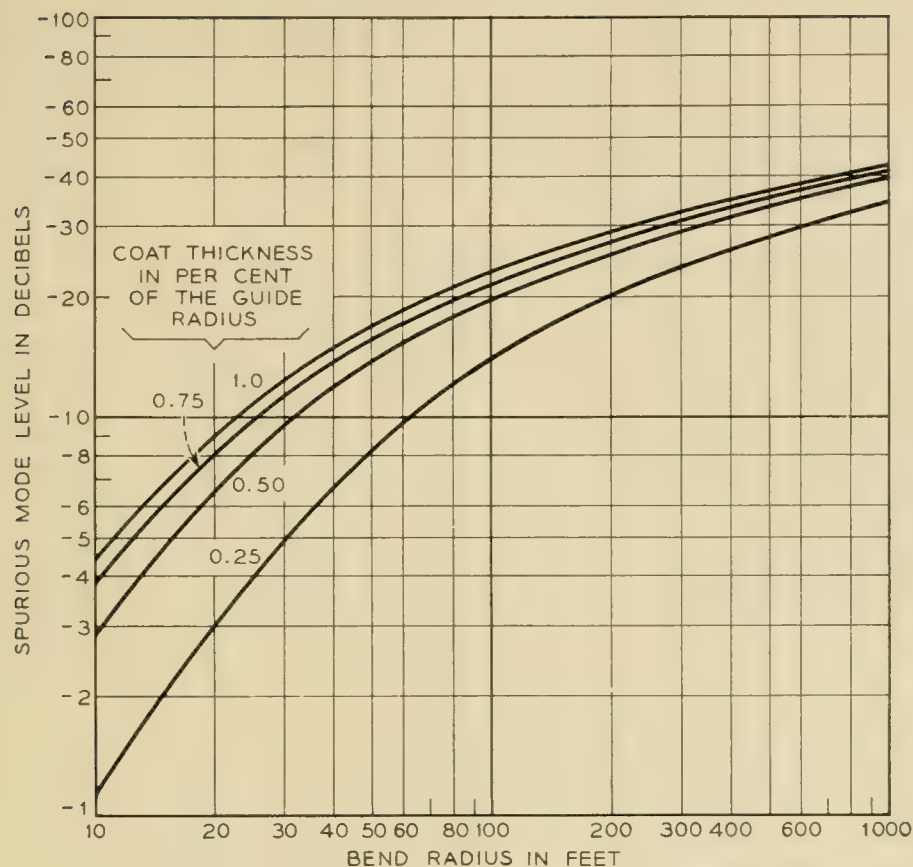


Fig. 5 — Spurious mode level in a uniform bend of a dielectric-coated guide. $2a = 2$ inches; $\epsilon = 2.5$; $\lambda = 5.4$ mm.

in the normal mode. The bending radius R is a function of position and an average bending radius,

$$\frac{1}{R_{Av}^2} = \frac{1}{z} \int_0^z \frac{dz}{R^2}, \quad (29)$$

has to be used in (20).

All the attenuation terms (20) decrease with increasing coat thickness δ ; the TE_{01} attenuation in the straight guide (10) and (14) increases with δ . Consequently there is an optimum thickness for which the total increase in attenuation is a minimum. This optimum thickness depends on the average radius of curvature. In Fig. 7 optimum thickness and the corresponding increase in attenuation have been plotted versus the average radius of curvature.

In gentle curvatures the normal mode is a mixture of TE_{01} and TM_{11} only and the attenuation increase as caused by the curvature is given by (27). In this case the optimum thickness and the corresponding attenuation increase are:

$$\delta_{\text{opt}} = \frac{1}{\sqrt[4]{2}} \frac{1}{p_{01}} \frac{\sqrt{\epsilon'}}{(\epsilon' - 1)^{3/4}} \sqrt{\frac{a}{R_{\text{Av}}}}, \quad (30)$$

$$\frac{\Delta\alpha}{\alpha_{01}} = \frac{\sqrt{2}}{\nu_{01}^2} \frac{\epsilon'}{\sqrt{\epsilon' - 1}} \frac{a}{R_{\text{Av}}}. \quad (31)$$

We note that δ_{opt} does not depend on frequency.

So far we have listed only the useful properties of the dielectric-coated guide. There is, however, one serious disadvantage. Serpentine bends caused by equally spaced discrete supports and the elasticity of the pipe are inherently present in any waveguide line. At critical frequencies, when the supporting distance l is a multiple of the beat wavelength λ_b between TE_{01} and a particular coupled mode, an increase in TE_{01} attenuation (23) and a spurious mode level (24) result from the mode conversion.

We evaluate (23) and (24) for a dielectric-coated copper pipe of 2.000-inch I.D. and 2.375-inch O.D. and a supporting distance $l = 15$ ft. The

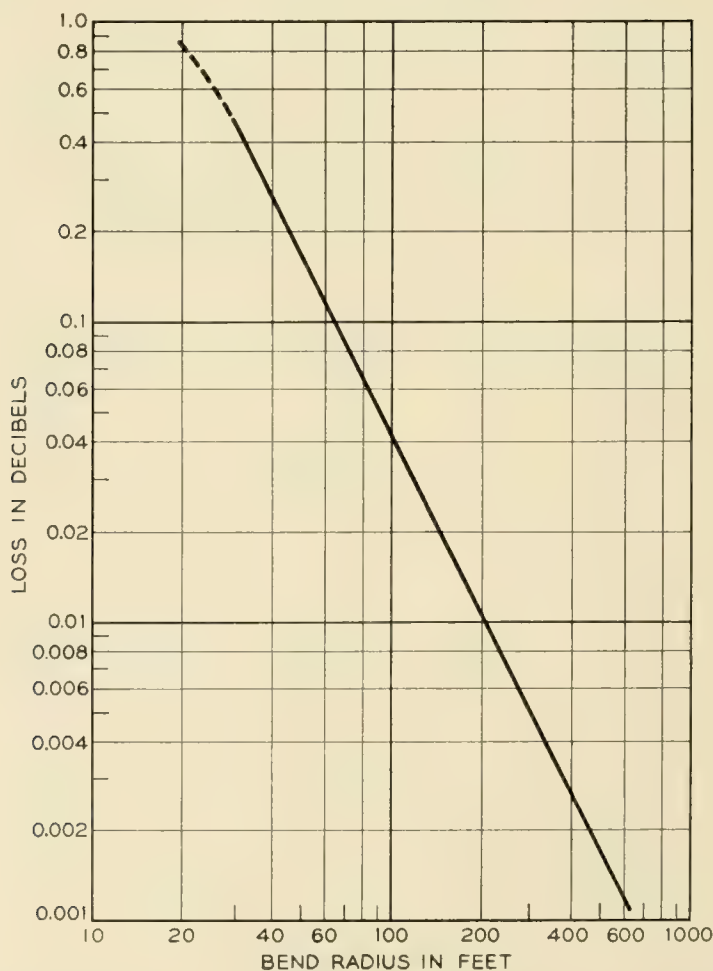


Fig. 6 — Bend conversion loss in a dielectric-coated guide of optimum design according to equation (28). $2a = 2$ inches; $\delta = 1.25\%$; $\epsilon = 2.5$; $\lambda = 5.4$ mm.

result is for coat thicknesses which cause the wavelength $\lambda = 5.4$ mm to be critical with respect to TE_{01} - TM_{11} coupling:

$$\begin{aligned} \delta = 0.002 \quad l = \lambda_b \quad \frac{\Delta\alpha_s}{\alpha_{01}} = 43.3 \quad 20 \log_{10} \left| \frac{E_2}{E_1} \right| &= -1.29 \text{ db} \\ &= 0.004 \quad = 2\lambda_b \quad = 2.70 \quad = -13.32 \text{ db} \\ &= 0.006 \quad = 3\lambda_b \quad = 0.169 \quad = -25.37 \text{ db.} \end{aligned}$$

The values corresponding to $\delta = 0.002$ do not satisfy the condition (25). Therefore they cannot be considered as a quantitative result but only as an indication that the mode conversion is very high. We conclude from these values that the mode conversion is much too high for a coat thickness which makes the beat wavelength between TE_{01} and TM_{11} equal to the supporting distance or half of it.

If no other measures can be taken, such as removing the periodicity of the supports or inserting mode filters, the lowest critical frequencies of TE_{01} - TM_{11} coupling have to be avoided. A dielectric coat must be

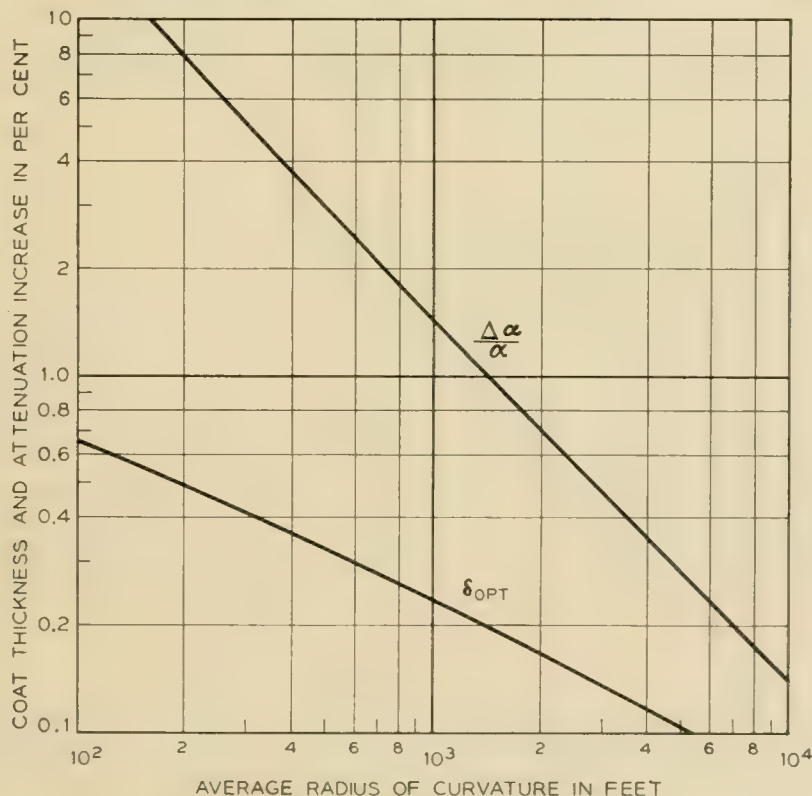


Fig. 7 — Optimum coat thickness and increase in TE_{01} attenuation of a dielectric-coated guide with random deviations from straightness. $2a = 2$ inches; $\lambda = 5.4$ mm; $\epsilon = 2.5$.

chosen, which is thin enough or thick enough to keep these critical frequencies out of the band.

Mode conversion from TE_{01} to the TE_{1m} modes in serpentine bends is not changed substantially by the presence of the dielectric coat. Generally these mode conversion effects decrease rapidly with the beat wavelength. When the curvature varies slowly compared to the difference in phase constant between coupled waves almost pure normal mode propagation is maintained.⁷

V. MODE CONVERSION AT TRANSITIONS FROM PLAIN WAVEGUIDE TO THE DIELECTRIC-COATED WAVEGUIDE

We consider a round waveguide, a section of which has a dielectric layer next to the walls. A pure TE_{01} wave incident on this dielectric-coated section will excite TE_{0m} waves. For circular electric wave transmission it is important to keep low the power level of the higher circular electric waves which have low loss.

In an evaluation of Schelkunoff's generalized telegraphist's equations for the TE_{01} mode in a circular waveguide containing an inhomogeneous dielectric¹ S. P. Morgan describes the wave propagation in the dielectric-filled waveguide in terms of normal modes of the unfilled waveguide. The only restriction made in this analysis for the dielectric insert is

$$\frac{1}{S} \int_S |\epsilon - 1| dS \ll 1, \quad (32)$$

where S is the cross-sectional area of the guide.

The dielectric-coated guide satisfies (32), and we may use the results of Morgan's evaluation here.

The round waveguide is considered as an infinite set of transmission lines, each of which represents a normal mode. Along the dielectric-coated section the TE_{0m} transmission lines are coupled mutually. The coupling coefficient d between TE_{01} and one of the TE_{0m} waves is obtained by taking Morgan's general formula and evaluating it for the dielectric coat:

$$d = p_{01}p_{0m} \frac{(\epsilon' - 1)\beta^2}{\sqrt{\beta_{01}\beta_{0m}}} \frac{\delta^3}{3}. \quad (33)$$

We introduce this coupling coefficient into the coupled line equations (16). Since $d/\Delta\beta \ll 1$ for any of the coupled modes, the spurious mode level at the output of the dielectric-coated section is given by (21),

$$20 \log_{10} \left| \frac{E_{2\max}}{E_1} \right| = 20 \log_{10} \frac{2}{3} \frac{(\epsilon' - 1)p_{01}p_{0m}\beta^2}{(\beta_{01} - \beta_{0m})\sqrt{\beta_{01}\beta_{0m}}} \delta^3. \quad (34)$$

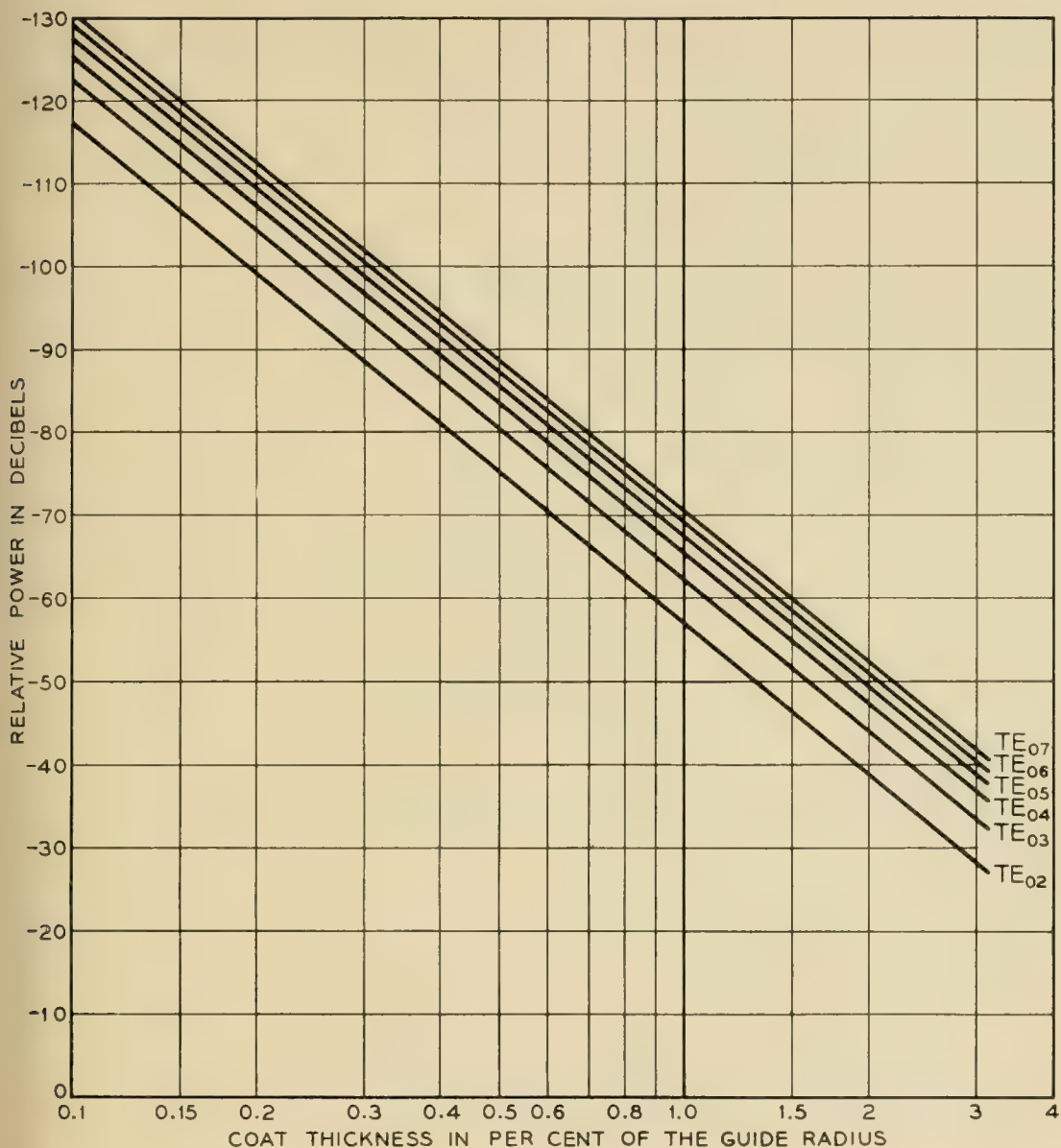


Fig. 8 — Higher circular electric waves at a transition from plain waveguide to the dielectric-coated guide. $\epsilon = 2.5$; $a/\lambda = 4.70$.

Fig. 8 shows an evaluation of (34) for $a/\lambda = 4.70$ and $\epsilon' = 2.5$ corresponding to a 2-inch pipe with a polystyrene coat at $\lambda = 5.4$ mm.

VI. SUMMARY

A theoretical analysis of wave propagation in the dielectric-coated guide is presented to provide information necessary for circular electric wave transmission in this waveguide structure. The normal modes of a waveguide with a thin dielectric coat are perturbed modes of the plain waveguide. While the perturbation of the phase constant is only of third order of the coat thickness for the circular electric waves, it is of first

order for all other modes. Thus, the degeneracy of equal phase constants of circular electric waves and TM_{1m} waves can be removed quite effectively. The additional attenuation to TE_{01} wave as caused by the dielectric loss in the coat and the increased wall current loss remains small as long as the coat is thin.

The dielectric-coated guide may be used for negotiating intentional bends or for avoiding extreme straightness requirements on normally straight sections. For intentional bends of as small a radius as possible, an optimum thickness of the coat makes the mode-conversion losses to TM_{11} and TE_{12} modes equal and minimizes the total conversion loss. For random deviations from straightness, an average radius of curvature is defined. For this average radius an optimum thickness of the coat minimizes the additional TE_{01} attenuation as caused by curvature and dielectric coat. In random deviations from straightness, propagation of only one normal mode is maintained as long as only the rate of change of curvature is small compared to the square of the difference in phase constant between TE_{01} and any coupled mode.

Serpentine bends, caused by equally spaced supports and the associated elastic deformation of the pipe, increase the TE_{01} attenuation substantially at certain critical frequencies, when the supporting distance is a multiple of the beat wavelength. The lowest critical frequencies of TE_{01} - TM_{11} coupling corresponding to a beat wavelength, which is equal to the supporting distance or half of it, have to be avoided by choosing the proper coat thickness.

At transitions from plain waveguide to dielectric-coated guide higher circular electric waves are excited by the TE_{01} wave. However, the power level of these spurious modes is low for a thin dielectric coat.

ACKNOWLEDGMENTS

The dielectric-coated waveguide is the subject of two patents.⁹ Some of its useful properties were brought to the writer's attention by a communication between the Standard Telecommunication Laboratory, Ltd., England, and S. E. Miller. For helpful discussions the writer is indebted to E. A. J. Marcatili, S. E. Miller, and D. H. Ring.

APPENDIX I

Approximate Solutions of the Characteristic Equation

In the following calculation we will use the definitions:

$$\begin{aligned} R_n(x, \rho x) &= J_n(x) N_{n-1}(\rho x) - J_{n-1}(\rho x) N_n(x), \\ S_n(x, \rho x) &= J_{n-1}(x) N_n(\rho x) - J_n(\rho x) N_{n-1}(x), \end{aligned} \tag{35}$$

and their relation to the definitions (2);

$$\begin{aligned}\frac{W_n(x, \rho x)}{U_n(x, \rho x)} &\equiv \rho x \frac{R_n(x, \rho x)}{U_n(x, \rho x)} + n \equiv -\rho x \frac{S_{n+1}(x, \rho x)}{U_n(x, \rho x)} - n, \\ \frac{V_n(x, \rho x)}{Z_n(x, \rho x)} &\equiv \rho x \frac{x U_{n-1}(x, \rho x) + n R_n(x, \rho x)}{x S_n(x, \rho x) + n U_n(x, \rho x)} + n.\end{aligned}\quad (36)$$

To find solutions of (1) for $1 - \rho \equiv \delta \ll 1$ we first substitute for the functions (36) their series expansions with respect to δ .

We have for instance

$$U_n(x, \rho x) = N_n(x) J_n(x - \delta x) - J_n(x) N_n(x - \delta x)$$

where we introduce

$$\begin{aligned}J_n(x - \delta x) &= J_n(x) + \delta x J_n'(x) + \frac{(\delta x)^2}{2} J_n''(x) + \dots \\ &= J_n(x) + \delta x \left[J_{n+1}(x) - \frac{n}{x} J_n(x) \right] \\ &\quad + \frac{(\delta x)^2}{2} \left[\frac{1}{x} J_{n+1}(x) + \left(\frac{n^2 - n}{x^2} - 1 \right) J_n(x) \right] + \dots,\end{aligned}$$

and

$$\begin{aligned}N_n(x - \delta x) &= N_n(x) + \delta x \left[N_{n+1}(x) - \frac{n}{x} N_n(x) \right] \\ &\quad + \frac{(\delta x)^2}{2} \left[\frac{1}{x} N_{n+1}(x) + \left(\frac{n^2 - n}{x^2} - 1 \right) N_n(x) \right] + \dots.\end{aligned}$$

Upon using the relation

$$J_{n+1}(x) N_n(x) - J_n(x) N_{n+1}(x) = \frac{2}{\pi x}$$

we get

$$U_n(x, \rho x) = \frac{2}{\pi} \delta \left[1 + \frac{\delta}{2} + \frac{\delta^2 x^2}{6} \left(\frac{n^2 + 2}{x^2} - 1 \right) \right] + O(\delta^4), \quad (37)$$

and by the same procedure

$$\begin{aligned}R_n(x, \rho x) &= \frac{2}{\pi x} \left[1 - (n-1)\delta + \left(\frac{(n-1)(n-2)}{x^2} - 1 \right) \frac{(\delta x)^2}{2} \right. \\ &\quad \left. - \frac{n-2}{x} \left(1 + \frac{(n-1)(n-3)}{x^2} \right) \frac{(\delta x)^3}{6} \right] + O(\delta^4).\end{aligned}\quad (38)$$

With these expressions the functions (36) are approximated by:

$$\begin{aligned}\frac{W_n(x, \rho x)}{U_n(x, \rho x)} &= \frac{1}{\delta} \left(1 - \frac{\delta}{2}\right) + 0(\delta), \\ \frac{V_n(x, \rho x)}{Z_n(x, \rho x)} &= \delta(x^2 - n^2) + 0(\delta^2),\end{aligned}\quad (39)$$

and for $n = 0$ especially:

$$\begin{aligned}\frac{1}{\rho x^2} \frac{V_0(x, \rho x)}{Z_0(x, \rho x)} &= -\frac{1}{x} \frac{U_1(x, \rho x)}{R_1(x, \rho x)} \\ &= -\delta \left[1 + \frac{\delta}{2} + \delta^2 \left(\frac{1}{3}x^2 + \frac{1}{2}\right)\right] + 0(\delta^4).\end{aligned}\quad (40)$$

Since for $\delta = 0$ the roots of $J_n(\rho x_1) = 0$ and $J'_n(\rho x_1) = 0$ represent the solutions of (1) for the TM and TE waves respectively, we expand the Bessel functions of the argument ρx_1 in series around these roots:

$$\rho x_1 = (1 - \delta)(p_{nm} + \Delta x).$$

The result for TM waves is

$$J_n(p_{nm}) \equiv 0: \quad \frac{J'_n(\rho x_1)}{J_n(\rho x_1)} = \frac{1}{\Delta x - \delta p_{nm}}; \quad (41)$$

for TE waves:

$$J'_n(p_{nm}) \equiv 0: \quad \frac{J_n(\rho x_1)}{J'_n(\rho x_1)} = \frac{n^2 - p_{nm}^2}{p_{nm}^2} (x - \delta p_{nm});$$

and for TE_{0m} waves especially:

$$\begin{aligned}J'_0(p_{0m}) &= 0: \\ -\frac{J'_0(\rho x_1)}{J_0(\rho x_1)} &= \Delta x - \delta p_{0m} - \frac{1}{2p_{0m}} (\Delta x^2 + \delta^2 p_{0m}^2) - \frac{\delta^3}{6} p_{0m} (3 + 2p_{0m}^2).\end{aligned}\quad (42)$$

Introducing (39 ... 42) into the characteristic equation (1) and neglecting higher order terms in δ and Δx we get the following approximations for (1):

$$\begin{aligned}TM_{nm} \text{ waves:} \quad \Delta x &= \left(p_{nm} - \frac{1}{\epsilon} \frac{x_2^2}{p_{nm}}\right) \delta \\ TE_{nm} \text{ waves:} \quad \Delta x &= \frac{x_2^2 - p_{nm}^2}{p_{nm}^2 - n^2} \frac{n^2}{\epsilon p_{nm}} \delta\end{aligned}\quad (43)$$

$$TE_{om} \text{ waves:} \quad \Delta x = -\frac{p_{om}}{3} (x_2^2 - p_{om}^2) \delta^3. \quad (44)$$

If we write for the perturbed propagation constant

$$\gamma = \gamma_{nm} + \Delta\gamma,$$

we get with (3) and (4):

$$\begin{aligned} \Delta x &= a^2 \frac{\gamma_{nm}}{p_{nm}} \Delta\gamma \\ &= ja \frac{\sqrt{1 - \nu_{nm}^2}}{\nu_{nm}} \Delta\gamma, \end{aligned}$$

and

$$x_2^2 = (\epsilon - 1 + \nu_{nm}^2) \frac{p_{nm}^2}{\nu_{nm}^2},$$

in which

$$\nu_{nm} = \frac{\lambda}{\lambda_{cnm}} = \frac{p_{nm}}{2\pi} \frac{\lambda}{a}.$$

Using these expressions in (43) and (44) we finally get approximate formulae for the perturbation of the propagation constant as caused by the dielectric coat:

$$\begin{aligned} TM_{nm} \text{ waves:} \quad \frac{\Delta\gamma}{\gamma_{nm}} &= \frac{\epsilon - 1}{\epsilon} \delta \\ TE_{nm} \text{ waves:} \quad \frac{\Delta\gamma}{\gamma_{nm}} &= \frac{n^2}{p_{nm}^2 - n^2} \frac{\epsilon - 1}{\epsilon} \frac{1}{1 - \nu_{nm}^2} \delta \quad (45) \\ TE_{0m} \text{ waves:} \quad \frac{\Delta\gamma}{\gamma_{0m}} &= \frac{p_{0m}^2}{3} \frac{\epsilon - 1}{1 - \nu_{0m}^2} \delta^3. \end{aligned}$$

The series expansions used so far hold only when $(1 - \rho)x \ll 1$. Approximate expressions which require only $(1 - \rho) \ll 1$ are:⁸

$$\begin{aligned} U_n(x, \delta x) &= \frac{2}{\pi} \frac{\sin \frac{1 - \rho}{\sqrt{\rho}} \sqrt{\rho x^2 - n^2}}{\sqrt{\rho x^2 - n^2}}, \\ S_n(x, \delta x) &= \frac{2}{\pi} \frac{1}{(1 - \rho)x} \left[\sqrt{\rho} \cos \left(\frac{1 - \rho}{\sqrt{\rho}} \sqrt{\rho x^2 - (n - 1)^2} \right) \right. \\ &\quad \left. - \frac{1}{\sqrt{\rho}} \cos \left(\frac{1 - \rho}{\sqrt{\rho}} \sqrt{\rho x^2 - n^2} \right) \right]. \end{aligned}$$

If $x^2 \gg n^2$ these expressions may be further simplified:

$$U_n(x, \rho x) = \frac{2}{\pi} \frac{1}{\sqrt{\rho x}} \sin (1 - \rho)x,$$

$$S_n(x, \rho x) = -\frac{2}{\pi} \frac{1}{\sqrt{\rho x}} \cos (1 - \rho)x.$$

Substituting these values in (36) we get:

$$\frac{W_n(x, \rho x)}{U_n(x, \rho x)} = \rho x \cot (1 - \rho)x - n,$$

$$\frac{V_n(x, \rho x)}{Z_n(x, \rho x)} = -\rho x \tan (1 - \rho)x \frac{1 - \frac{n}{x} \frac{1 - \rho}{\rho} \cot (1 - \rho)x}{1 - \frac{n}{x} \tan (1 - \rho)x}.$$

With the restrictions $(1 - \rho)x < \pi/2$ and $n \ll x$ we may write:

$$\frac{W_n(x, \rho x)}{U_n(x, \rho x)} = \rho x \cot (1 - \rho)x,$$

$$\frac{V_n(x, \rho x)}{Z_n(x, \rho x)} = -\rho x \tan (1 - \rho)x.$$
(46)

With (46) the characteristic equation (1) is

$$n^2 \left[\frac{1}{x_1^2} - \frac{1}{x_2^2} \right]^2 - \rho^2 \frac{x_2^2 - x_1^2}{x_2^2 - \epsilon x_1^2} \left[\frac{1}{x_1} \frac{J_n'(\rho x_1)}{J_n(\rho x_1)} + \frac{\epsilon}{x_2} \cot \delta x_2 \right] \cdot \left[\frac{1}{x_2} \frac{J_n'(\rho x_1)}{J_n(\rho x_1)} - \frac{1}{x_2} \tan \delta x_2 \right] = 0.$$

Solving this equation for $\cot \delta x_2$ we get:

$$\cot \delta x_2 = \frac{F}{2} + \sqrt{1 + \frac{F^2}{4}},$$
(47)

in which

$$F = \frac{x_1}{x_2} \frac{J_n(\rho x_1)}{J_n'(\rho x_1)} \left[1 + \frac{n^2}{\epsilon \rho^2} \frac{(x_1^2 - x_2^2)(x_2^2 - \epsilon x_1^2)}{x_1^4 x_2^2} \right] - \frac{x_2}{\epsilon x_1} \frac{J_n'(\rho x_1)}{J_n(\rho x_1)}.$$

To evaluate (47) we specify a value of $\Delta\beta/\beta_{nm}$ and enter the right hand side of (47) with

$$x_1^2 = p_{nm}^2 - 2 \frac{\Delta\beta}{\beta_{nm}} \beta_{nm}^2 a^2 \left(1 + \frac{1}{2} \frac{\Delta\beta}{\beta_{nm}} \right),$$

$$x_2^2 = (\epsilon - 1) \beta_{nm}^2 a^2 + x_1^2,$$

$$\rho = 1 - \delta_1$$

in which δ_1 is a first order approximation as given by (10) for a particular mode. Equation (47) then yields a value δ_2 which is usually accurate enough. To improve the accuracy the same calculation is repeated using δ_2 . Since for small values of δ a change in δ affects the right hand side only slightly this method converges rapidly.

APPENDIX II

The TE₀₁ Attenuation in the Dielectric-Coated Waveguide

The attenuation constant of a transmission line can be expressed as

$$\alpha = \frac{1}{2} \frac{P_M}{P}$$

in which P_M is the power dissipated per unit length in the line and P is the total transmitted power. For the dielectric-coated guide the power P_M is dissipated in the metal walls with finite conductivity σ .

For TE_{0m} waves the wall currents are $i = H_z(a)$ (H_z axial component of the magnetic field) and therefore

$$\begin{aligned} P_M &= \frac{1}{2} \int_0^{2\pi} \frac{i i^*}{t \cdot \sigma} a d\varphi \\ &= \frac{\pi a}{t \sigma} H_z(a) H_z^*(a) \end{aligned}$$

where t = skin depth and σ = conductivity.

The total transmitted power is

$$P = -\frac{1}{2} \int_0^{2\pi} \int_0^a E_\varphi H_r^* r dr d\varphi.$$

We introduce expressions for the field components, which are listed elsewhere^{3, 4} and carry out the integration. Finally we express the wall current attenuation in terms of the functions (2) and (35);

$$\begin{aligned} \alpha_M &= \frac{\alpha_{cm}}{p_{om}^2} x_2^2 \sqrt{\frac{x_2^2 - x_1^2}{x_2^2 - \epsilon x_1^2}} \left\{ 1 + \frac{\pi^2}{2} \rho x_2 \right. \\ &\quad \cdot \left(\frac{x_2^2}{x_1^2} - 1 \right) \left(U_1(x_2, \rho x_2) + \frac{\rho x_2}{2} R_1(x_2, \rho x_2) \right) R_1(x_2, \rho x_2) \left. \right\}^{-1} \end{aligned} \quad (48)$$

To get an approximation for a thin dielectric coat we use the expressions (37) and (38) and obtain

$$\frac{\Delta \alpha_M}{\alpha_{om}} = \frac{\alpha_M - \alpha_{cm}}{\alpha_{om}} = (\epsilon - 1) \frac{p_{om}^2}{v_{om}^2} \delta^2. \quad (49)$$

REFERENCES

1. S. P. Morgan, Theory of Curved Circular Waveguide Containing an Inhomogeneous Dielectric, pp. 1209-1251, this issue.
2. S. E. Miller, Coupled Wave Theory and Waveguide Applications, B.S.T.J., **33**, 661-719, May, 1954.
3. H. Buchholz, Der Hohlleiter von kreisförmigem Querschnitt mit geschichtetem dielektrischen Einsatz, Ann. Phys., **43**, pp. 313-368, 1943.
4. H. M. Wachowski and R. E. Beam, Shielded Dielectric Rod Waveguides, Final Report on Investigations of Multi-Mode Propagation in Waveguides and Microwave Optics, Microwave Laboratory, Northwestern University, Ill., 1950.
5. S. P. Morgan and J. A. Young, Helix Waveguide, B.S.T.J., **35**, pp. 1347-1384, November, 1956.
6. H. G. Unger, Circular Electric Wave Transmission through Serpentine Bends, pp. 1279-1291, this issue.
7. J. S. Cook, Tapered Velocity Couplers, B.S.T.J., **34**, pp. 807-822, July, 1955; A. G. Fox, Wave Coupling by Warped Normal Modes, B.S.T.J., **34**, pp. 823-852, July, 1955; W. H. Louisell, Analysis of the Single Tapered Mode Coupler, B.S.T.J., **34**, pp. 853-870, July, 1955.
8. H. Buchholz, Approximation Formulae for a Well Known Difference of Products of Two Cylinder Functions, Phil. Mag., Series 7, **27**, pp. 407-420, 1939.
9. A. E. Karbowski, British Patent No. 751,322. H. G. Unger and O. Zinke, German Patent No. 935,677.

Circular Electric Wave Transmission Through Serpentine Bends

By H. G. UNGER

(Manuscript received January 9, 1957)

An otherwise straight waveguide line with equally spaced discrete supports may deform elastically into a serpentine bend under its own weight. The TE_{01} wave couples in such bends to the TM_{11} and TE_{1n} waves. The general solution of coupled lines with varying coupling coefficient is applied to a serpentine bend by an iterative process, and evaluated for the elastic curve resulting from a periodically supported line. TE_{01} - TM_{11} coupling causes only a small increase in TE_{01} attenuation. Mode conversion to TE_{1n} waves can become seriously high at certain critical frequencies when the supporting distance is a multiple of the beat wavelength. In a copper pipe of $2\frac{3}{8}$ inch O.D. and 2 inch I.D., the mode conversion to the TE_{12} wave at critical frequencies near 5-mm wavelength causes a TE_{01} attenuation increase of 90 per cent and a spurious mode level of -7 db. These mode conversion effects can be controlled effectively by inserting mode filters.

I. INTRODUCTION

In curved sections of round waveguide the TE_{01} - wave couples to the TM_{11} and TE_{1n} waves, and power is converted to these waves when the TE_{01} wave is transmitted through bends. A form of bend which is inherently present even in an otherwise straight and perfect line is the serpentine bend, Fig. 1. Between discrete supports the pipe is deflected by the force of its own weight. The resulting curve is well known from the theory of elasticity. The curvature varies along the axis following essentially a square law. The minimum bending radius occurs at the supports. For the practical example of a copper pipe of $2\frac{3}{8}$ -inch O.D., 2.00-inch I.D. and a supporting distance of 15 ft, this minimum bending radius is 992 ft. A uniform bend of this radius would convert most of the power incident in the TE_{01} wave to the TM_{11} wave after a certain length of bend.

Fortunately a serpentine bend with the same order of bending radius does not affect circular electric wave transmission as seriously as does a

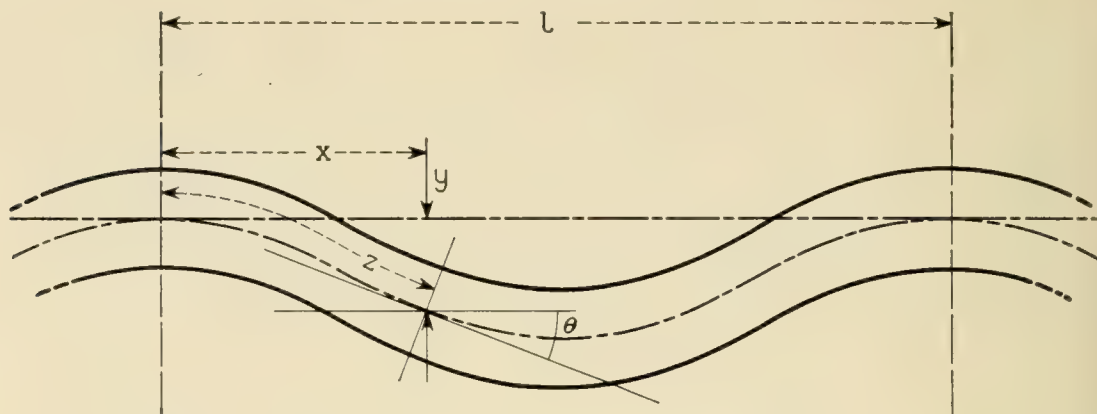


Fig. 1 — Serpentine bends.

uniform bend with the same bending radius. This will be shown in the following analysis.

A treatment of serpentine bends given previously by Albersheim¹ considers only circular and sinusoidal curvatures. Furthermore, it does not show all the effects of serpentine bends that we are interested in. We shall present a more general and complete analysis here. The only restriction we have to make is that the power exchanged in one section of the serpentine bend from the TE_{01} mode to any of the coupled modes or vice versa is small compared to the power in the mode from which it has been abstracted. The curvature may be any function of distance along the serpentine bend.

The general results indicate that normally a serpentine bend causes an additional attenuation to the TE_{01} mode. Part of the TE_{01} power which travels temporarily in one of the coupled modes suffers the higher attenuation of this coupled mode. Formulae for this increase in attenuation constant are obtained for periodically supported guides from the deflection curve given by the theory of elasticity.

The coupling between the TE_{01} and TM_{11} waves causes only a very slight increase in TE_{01} attenuation, since the difference in propagation constant for these two modes is very small. In fact, if there were no difference in propagation constant, as in the round guide of infinite wall conductivity, coupling between TE_{01} and TM_{11} waves in serpentine bends would not affect the TE_{01} transmission at all.

Coupled modes which have a larger difference in phase constant than TM_{11} but still are close to TE_{01} cause a serious increase in TE_{01} attenuation at certain critical frequencies. If a multiple of the beat wavelength between the TE_{01} mode and a particular coupled mode is equal to the supporting distance, power is converted continuously into the coupled

mode. For low attenuation in the coupled mode the power transfer may even be essentially complete.

II. SOLUTION OF THE COUPLED LINE EQUATIONS

In a curved round waveguide the TE_{01} mode couples to the TM_{11} mode and the infinite set of TE_{1n} waves. Consequently, the wave propagation is described by an infinite system of simultaneous first order linear differential equations. An adequate procedure is to consider only coupling between TE_{01} and one of the spurious modes at a time. Thus, the infinite system of equations reduces to the well known coupled line equations,²

$$\begin{aligned}\frac{dE_1}{dz} + \gamma_1 E_1 - k E_2 &= 0, \\ \frac{dE_2}{dz} + \gamma_2 E_2 - k E_1 &= 0;\end{aligned}\tag{1}$$

in which

$E_{1,2}(z)$ = wave amplitudes in mode 1 (here always TE_{01}) and mode 2 (TM_{11} or one of the TE_{1n}), respectively;

$\gamma_{1,2}$ = propagation constant of modes 1 and 2, respectively, (The small perturbation of γ_1 , γ_2 caused by the coupling may be neglected here); and

$k(z) = jc(z)$ = coupling coefficient between modes 1 and 2.

In the curved waveguide the coupling is proportional to the curvature;

$$c = \frac{c_0}{R} = c_0 \frac{d\theta}{dz},\tag{2}$$

in which R = radius of curvature, and θ = direction of guide axis. (The various coupling coefficients are listed in the appendix.) Without loss of generality we start with $\theta(0) = 0$. We will use the average propagation constant γ ,

$$\begin{aligned}\gamma &= \frac{1}{2}(\gamma_1 + \gamma_2), \\ \Delta\gamma &= \frac{1}{2}(\gamma_1 - \gamma_2).\end{aligned}\tag{3}$$

Several coordinate transformations will change (1) to a form which can be solved approximately. A similar procedure has been used to solve the

related problem of the tapered mode coupler.³ With the first transformation

$$\begin{aligned} E_1(z) &= \frac{1}{2}e^{-\gamma z} [u_1(z)e^{jc_0\theta} + u_2(z)e^{-jc_0\theta}], \\ E_2(z) &= \frac{1}{2}e^{-\gamma z} [u_1(z)e^{jc_0\theta} - u_2(z)e^{-jc_0\theta}], \end{aligned} \quad (4)$$

the new coordinates satisfy the equations:

$$\begin{aligned} \frac{du_1}{dz} + \Delta\gamma e^{-j2c_0\theta} u_2 &= 0, \\ \frac{du_2}{dz} + \Delta\gamma e^{j2c_0\theta} u_1 &= 0. \end{aligned} \quad (5)$$

With the second transformation

$$\begin{aligned} u_1(z) &= \frac{1}{2}[v_1(z)e^{-\Delta\gamma z} + v_2(z)e^{+\Delta\gamma z}], \\ u_2(z) &= \frac{1}{2}[v_1(z)e^{-\Delta\gamma z} - v_2(z)e^{+\Delta\gamma z}], \end{aligned} \quad (6)$$

$v_1(z)$ and $v_2(z)$ satisfy

$$\begin{aligned} \frac{dv_1}{dz} - 2\Delta\gamma \sin^2 c_0\theta v_1 + j\Delta\gamma \sin 2c_0\theta e^{2\Delta\gamma z} v_2 &= 0, \\ \frac{dv_2}{dz} + 2\Delta\gamma \sin^2 c_0\theta v_2 - j\Delta\gamma \sin 2c_0\theta e^{2\Delta\gamma z} v_1 &= 0. \end{aligned} \quad (7)$$

With the third transformation

$$\begin{aligned} v_1(z) &= w_1(z) \exp \left(2\Delta\gamma \int_0^z \sin^2 c_0\theta dz' \right), \\ v_2(z) &= w_2(z) \exp \left(-2\Delta\gamma \int_0^z \sin^2 c_0\theta dz' \right). \end{aligned} \quad (8)$$

We finally get the equations

$$\begin{aligned} \frac{dw_1}{dz} + \xi_1(z)w_2 &= 0, \\ \frac{dw_2}{dz} + \xi_2(z)w_1 &= 0, \end{aligned} \quad (9)$$

in which

$$\xi_{1,2}(z) = \pm j\Delta\gamma \sin 2c_0\theta \exp \left[\pm 2\Delta\gamma \left(z - 2 \int_0^z \sin^2 c_0\theta \, dz' \right) \right]. \quad (10)$$

As long as we have the condition $\int_0^z |\xi_{1,2}(z')| \, dz' \ll 1$, approximate solutions of (9) can be written down which proceed essentially in powers of $\xi_{1,2}$ as follows,

$$\begin{aligned} w_1(z) &= w_1(0) - w_2(0) \int_0^z \xi_1(z') \, dz' \\ &\quad + w_1(0) \int_0^z \xi_1(z') \int_0^{z'} \xi_2(z'') \, dz'' \, dz', \\ w_2(z) &= w_2(0) - w_1(0) \int_0^z \xi_2(z') \, dz' \\ &\quad + w_2(0) \int_0^z \xi_2(z') \int_0^{z'} \xi_1(z'') \, dz'' \, dz'. \end{aligned} \quad (11)$$

The new coordinates are related to the wave amplitudes by:

$$\begin{aligned} E_1(z) &= \frac{1}{2} e^{-\gamma z} \left\{ w_1(z) \exp \left[-\Delta\gamma \left(z - 2 \int_0^z \sin^2 c_0\theta \, dz' \right) \right] \cos c_0\theta \right. \\ &\quad \left. + jw_2(z) \exp \left[\Delta\gamma \left(z - 2 \int_0^z \sin^2 c_0\theta \, dz' \right) \right] \sin c_0\theta \right\}, \\ E_2(z) &= \frac{1}{2} e^{-\gamma z} \left\{ jw_1(z) \exp \left[-\Delta\gamma \left(z - 2 \int_0^z \sin^2 c_0\theta \, dz' \right) \right] \sin c_0\theta \right. \\ &\quad \left. + w_2(z) \exp \left[\Delta\gamma \left(z - 2 \int_0^z \sin^2 c_0\theta \, dz' \right) \right] \cos c_0\theta \right\}. \end{aligned} \quad (12)$$

The solution (12) in combination with (11) is general and may be applied to any form of curvature as long as the converted power remains small compared to the original power in either of the modes.

III. WAVE PROPAGATION IN SERPENTINE BENDS

If we apply (12) to a section of a serpentine bend with the length l we have $\theta(l) = 0$. The output amplitudes are related to the input amplitudes of both modes by a transmission matrix

$$E_i(l) = \| T \| E_i(0). \quad (13)$$

The elements of $\| T \|$ are obtained from (11) and (12):

$$\begin{aligned}
T_{11} &= \left[1 + \int_0^l \xi_1(z) \int_0^z \xi_2(z') dz' dz \right] \\
&\quad \cdot \exp \left[-\gamma_1 l + 2\Delta\gamma \int_0^l \sin^2 c_0 \theta dz \right], \\
T_{21} &= -\int_0^l \xi_1(z) dz \exp \left[-\gamma_1 l + 2\Delta\gamma \int_0^l \sin^2 c_0 \theta dz \right], \\
T_{12} &= -\int_0^l \xi_2(z) dz \exp \left[-\gamma_2 l - 2\Delta\gamma \int_0^l \sin^2 c_0 \theta dz \right], \\
T_{22} &= \left[1 + \int_0^l \xi_2(z) \int_0^z \xi_1(z') dz' dz \right] \\
&\quad \cdot \exp \left[-\gamma_2 l - 2\Delta\gamma \int_0^l \sin^2 c_0 \theta dz \right].
\end{aligned} \tag{14}$$

For a line of iterative serpentine bends we apply the rules of matrix calculus. If $E_{10} = 1$ and $E_{20} = 0$ at the input of the first section, then the output amplitudes of the n^{th} section are:

$$\begin{aligned}
E_{1n} &= \frac{1}{2} \left[1 + \frac{T_{11} - T_{22}}{\sqrt{(T_{11} - T_{22})^2 + 4T_{12}T_{21}}} \right] e^{-ng_1} \\
&\quad + \frac{1}{2} \left[1 - \frac{T_{11} - T_{22}}{\sqrt{(T_{11} - T_{22})^2 + 4T_{12}T_{21}}} \right] e^{-ng_2}, \tag{15}
\end{aligned}$$

$$E_{2n} = \frac{T_{12}}{\sqrt{(T_{11} - T_{22})^2 + 4T_{12}T_{21}}} [e^{-ng_1} - e^{-ng_2}],$$

where

$$e^{-g_{1,2}} = \frac{1}{2} [T_{11} + T_{22} \pm \sqrt{(T_{11} - T_{22})^2 + 4T_{12}T_{21}}]. \tag{16}$$

Two limiting cases for the expressions (15) and (16) are of special interest:

$$1. \quad |T_{11} - T_{22}|^2 \gg 4 |T_{12} T_{21}|$$

$$\begin{aligned}
e^{-g_1} &= T_{11} + \frac{T_{12}T_{21}}{T_{11} - T_{22}} \\
E_{1n} &= e^{-ng_1} - \frac{T_{12}T_{21}}{(T_{11} - T_{22})^2} [e^{-ng_1} - e^{-ng_2}] \tag{17}
\end{aligned}$$

$$e^{-g_2} = T_{22} - \frac{T_{12}T_{21}}{T_{11} - T_{22}} \quad E_{2n} = \frac{T_{12}}{T_{11} - T_{22}} [e^{-ng_1} - e^{-ng_2}]$$

$$2. \quad 4 |T_{12} T_{21}| \gg |T_{11} - T_{22}|^2$$

$$\begin{aligned}
e^{-g_{1,2}} &= \frac{T_{11} + T_{22}}{2} \pm \sqrt{T_{12}T_{21}} \\
E_{1n} &= \frac{1}{2} (e^{-ng_1} + e^{-ng_2}) \\
E_{2n} &= \frac{1}{2} \sqrt{\frac{T_{12}}{T_{21}}} (e^{-ng_1} - e^{-ng_2}).
\end{aligned} \tag{18}$$

In case 1 the wave amplitude E_{1n} is only affected by a slight change of the propagation constant. The small additional term in the expression for E_{1n} can usually be neglected. The power in the wave E_{2n} is small compared to the E_{1n} power; T_{12} is usually of the same order of magnitude as T_{21} .

In case 2, however, a complete power transfer between E_{1n} and E_{2n} occurs cyclically if the loss is sufficiently low. Consequently, condition (18) has to be carefully avoided. Rather, condition (17) must always be satisfied.

IV. SERPENTINE BENDS FORMED BY ELASTIC CURVES

A section of the waveguide line between two supports deforms like a beam fixed at both ends. Under its own weight, w per unit length, such a beam will bend and form an elastic curve, Fig. 1, whose deflection from a straight line is given by:

$$y = \frac{wl^4}{24EI} \frac{x^2}{l^2} \left(1 - \frac{x}{l}\right)^2, \quad (19)$$

in which E = modulus of elasticity of the beam, and I = moment of inertia of the beam. Since we are concerned with small deflections y only, we have $x = z$ and $\theta = dy/dx$. Hence,

$$\theta = d \left(\frac{z}{l} - 3 \frac{z^2}{l^2} + 2 \frac{z^3}{l^3} \right), \quad (20)$$

in which $d = wl^3/12EI$. Introducing the elastic curve (20) into the transmission matrix (10) and (14) and performing the integrations with $\sin c_0\theta = c_0\theta$ we get for the elements of the transmission matrix:

$$\begin{aligned} T_{11} &= \exp \left[-\gamma_1 l + \Delta\gamma l \frac{c_0^2 d^2}{105} \right] \left\{ 1 + \frac{c_0^2 d^2}{4\Delta\gamma^6 l^6} \left[9 - 3\Delta\gamma^2 l^2 + \Delta\gamma^4 l^4 \right. \right. \\ &\quad \left. \left. + \frac{2}{5} \Delta\gamma^5 l^5 - \frac{4}{105} \Delta\gamma^7 l^7 - (3 - 3\Delta\gamma l + \Delta\gamma^2 l^2)^2 e^{2\Delta\gamma l} \right] \right\}, \\ T_{22} &= \exp \left[-\gamma_2 l - \Delta\gamma l \frac{c_0^2 d^2}{105} \right] \left\{ 1 + \frac{c_0^2 d^2}{4\Delta\gamma^6 l^6} \left[9 - 3\Delta\gamma^2 l^2 + \Delta\gamma^4 l^4 \right. \right. \\ &\quad \left. \left. - \frac{2}{5} \Delta\gamma^5 l^5 + \frac{4}{105} \Delta\gamma^7 l^7 - (3 + 3\Delta\gamma l + \Delta\gamma^2 l^2)^2 e^{-2\Delta\gamma l} \right] \right\}, \\ T_{12} &= T_{21} = j \frac{c_0 d}{2\Delta\gamma^3 l^3} [(3 - 3\Delta\gamma l + \Delta\gamma^2 l^2) e^{-\gamma_2 l} \\ &\quad - (3 + 3\Delta\gamma l + \Delta\gamma^2 l^2) e^{-\gamma_1 l}]. \end{aligned} \quad (21)$$

The expressions (21) are hard to evaluate, but for some special cases of interest they can be simplified greatly.

To compute the coupling effects between the TE_{01} wave and the TM_{11} wave we make use of $|\Delta\gamma l| \ll 1$ and get the following approximations:

$$\begin{aligned} T_{11} &= \left[1 - \frac{c_0^2 d^2}{315} \Delta\gamma^3 l^3 \right] \exp - \left(\gamma_1 - \frac{c_0^2 d^2}{105} \Delta\gamma \right) l, \\ T_{22} &= \left[1 + \frac{c_0^2 d^2}{315} \Delta\gamma^3 l^3 \right] \exp - \left(\gamma_2 + \frac{c_0^2 d^2}{105} \Delta\gamma \right) l, \\ T_{12} &= T_{21} = j \frac{c_0 d}{15} \Delta\gamma^2 l^2 e^{-\gamma l}. \end{aligned} \quad (22)$$

In (22) the condition $|T_{11} - T_{22}|^2 \gg 4 |T_{12} T_{21}|$ is satisfied; consequently the wave propagation is described by (17),

$$\begin{aligned} E_{1n} &= \exp - \left(\gamma_1 - \frac{c_0^2 d^2}{105} \Delta\gamma \right) nl + \frac{c_0^2 d^2}{450} \Delta\gamma^2 l^2 \sinh \Delta\gamma n l e^{-\gamma n l}, \\ E_{2n} &= j \frac{c_0 d}{15} \Delta\gamma l \sinh \Delta\gamma n l e^{-\gamma n l}. \end{aligned} \quad (23)$$

In addition to small oscillations, which are negligible, the wave amplitude E_1 suffers an additional attenuation

$$\Delta\alpha_s = -\frac{c_0^2 d^2}{105} \Delta\alpha. \quad (24)$$

Physically this means that to a first approximation there is no net power transfer from E_1 to E_2 . The power converted from E_1 to E_2 in one section of the iterative serpentine bends is all reconverted in the same section. But this power, which travels partly in the E_2 wave, suffers the E_2 attenuation and consequently changes the E_1 attenuation.

To evaluate (24) we introduce the coupling coefficient and the difference in attenuation constants between the TE_{01} and TM_{11} wave. Then, the relative increase of TE_{01} attenuation is:

$$\frac{\Delta\alpha_s}{\alpha_{01}} = 6.39 \times 10^{-3} d^2 \frac{a^2}{\lambda^2} \left(2.69 \frac{a^2}{\lambda^2} - 1 \right), \quad (25)$$

where α_{01} = attenuation constant of TE_{01} ,

a = inner radius of pipe,

λ = free-space wavelength.

A numerical example shows that the increase in TE_{01} attenuation caused by coupling to the TM_{11} wave in serpentine bends is small. For a copper pipe with 2.00-inch I.D. and $2\frac{3}{8}$ -inch O.D. and a supporting length $l = 15$ ft, we have $d = 1.51 \times 10^{-2}$, and at $\lambda = 5.4$ mm we get $\Delta\alpha_s/\alpha_{01} = 0.19 \times 10^{-2}$.

For coupling between the TE_{01} wave and the waves of the TE_{1n} family the difference in propagation constant is no longer small; the approximation which was valid for the TM_{11} wave can not be made for TE_{1n} waves. Actually, the supporting distance is usually several beat wavelengths. Therefore, no essential simplifications of (21) are possible for the general case of coupling between TE_{01} and TE_{1n} . But closer examination of (21) shows that if

$$2\Delta\beta l = 2m\pi \quad m = 1, 2, 3, \dots \quad (26)$$

is satisfied, the difference $T_{11}-T_{22}$ becomes very small. The net power converted to any of the TE_{1n} modes may be small in each section of the iterative serpentine bends, but if (26) is satisfied the contributions from each section add in phase and in a long line with the square of distance more and more power is built up in the particular TE_{1n} wave. Only when the attenuation in this TE_{1n} wave is large enough to damp out the power as fast as it is converted will an undistorted TE_{01} propagation be maintained. This condition for the attenuation constant can be derived from $|T_{11} - T_{22}|^2 \gg 4|T_{12}T_{21}|$. If $|\Delta\beta| \gg |\Delta\alpha|$ and $|\Delta\alpha l| \gg c_0^2 d^2 / 10m\pi$ then $T_{11} - T_{22} = -2\Delta\alpha l e^{-\gamma_1 l}$, and since $T_{12}T_{21} = -9 c_0^2 d^2 / m^4 \pi^4 e^{-2\gamma_1 l}$ the condition for undistorted TE_{01} propagation is:

$$|\Delta\alpha l|^2 \gg 9 \frac{c_0^2 d^2}{m^4 \pi^4}. \quad (27)$$

If (27) is satisfied the wave propagation is again described by (17). Neglecting all terms which are small because of (27) the wave amplitudes are:

$$\begin{aligned} E_{1n} &= \left[1 + 9 \frac{c_0^2 d^2}{m^4 \pi^4} \frac{1}{2\Delta\alpha l} \right]^n e^{-\gamma_1 n l} + 9 \frac{c_0^2 d^2}{m^4 \pi^4} \frac{1}{4\Delta\alpha^2 l^2} [1 - e^{2\Delta\alpha n l}] e^{-\gamma_1 n l}, \\ E_{2n} &= -j3 \frac{c_0 d}{m^2 \pi^2} \frac{1}{2\Delta\alpha l} [1 - e^{2\Delta\alpha n l}] e^{-\gamma_1 n l}. \end{aligned} \quad (28)$$

For not too large values of n , the first term of E_{1n} may be written:

$$\left[1 + 9 \frac{c_0^2 d^2}{m^4 \pi^4} \frac{1}{2\Delta\alpha l} \right]^n e^{-\gamma_1 n l} = \exp \left[- \left(\gamma_1 - 9 \frac{c_0^2 d^2}{m^4 \pi^4} \frac{1}{2\Delta\alpha l^2} \right) n l \right].$$

The additional attenuation to the E_1 wave as caused by the continuous

power abstraction is seen from this expression to be

$$\Delta\alpha_s = -9 \frac{c_0^2 d^2}{m^4 \pi^4} \frac{1}{2\Delta\alpha l^2}. \quad (29)$$

The condition (27) may now be written

$$2\Delta\alpha_s \ll |\Delta\alpha|, \quad (30)$$

or, the rate of conversion loss has to be small compared to the difference between E_2 attenuation and E_1 attenuation.

When the wave is travelling through a large number of serpentine bends, power is built up gradually in the E_2 mode to a constant value,

$$\left| \frac{E_2}{E_1} \right| = 3 \frac{c_0 d}{m^2 \pi^2} \left| \frac{1}{2\Delta\alpha l} \right|. \quad (31)$$

Both the attenuation increase and the power level in spurious modes can seriously affect the E_1 transmission.

To evaluate (29) and (31) we rewrite them with $m\pi = \Delta\beta l$ and $d = wl^3/12EI$ as follows;

$$\frac{\Delta\alpha_s}{\alpha_{01}} = - \left[\frac{w}{EI} \frac{c_0}{(2\Delta\beta)^2 \alpha_{01}} \right]^2 \frac{\alpha_{01}}{2\Delta\alpha}, \quad (32)$$

$$\left| \frac{E_2}{E_1} \right| = \frac{w}{EI} \frac{c_0}{(2\Delta\beta)^2 \alpha_{01}} \left| \frac{\alpha_{01}}{2\Delta\alpha} \right|. \quad (33)$$

We note from (32) and (33) that the coupling effects cannot be controlled by changing the supporting distance. Only the number of critical frequencies in a given range decreases with decreasing supporting distance.

The previously cited numerical example of the 2 inch copper pipe yields the following values at the critical frequencies of the two lowest TE_{1n} waves near $\lambda = 5.4$ mm:

$$\begin{aligned} TE_{11} \frac{\Delta\alpha_s}{\alpha_{01}} &= 0.114 & 20 \log \left| \frac{E_2}{E_1} \right| &= -23.4 \text{ db}, \\ TE_{12} \frac{\Delta\alpha_s}{\alpha_{01}} &= 0.855 & 20 \log \left| \frac{E_2}{E_1} \right| &= -6.85 \text{ db}. \end{aligned}$$

The mode conversion, especially to the TE_{12} wave, causes a seriously high additional attenuation and spurious mode level.

V. MODE FILTERS IN SERPENTINE BENDS

Periodically spaced supports are a condition for the critical case described by (32) and (33). Accordingly, the coupling effects can be controlled by removing the periodicity of the supports.

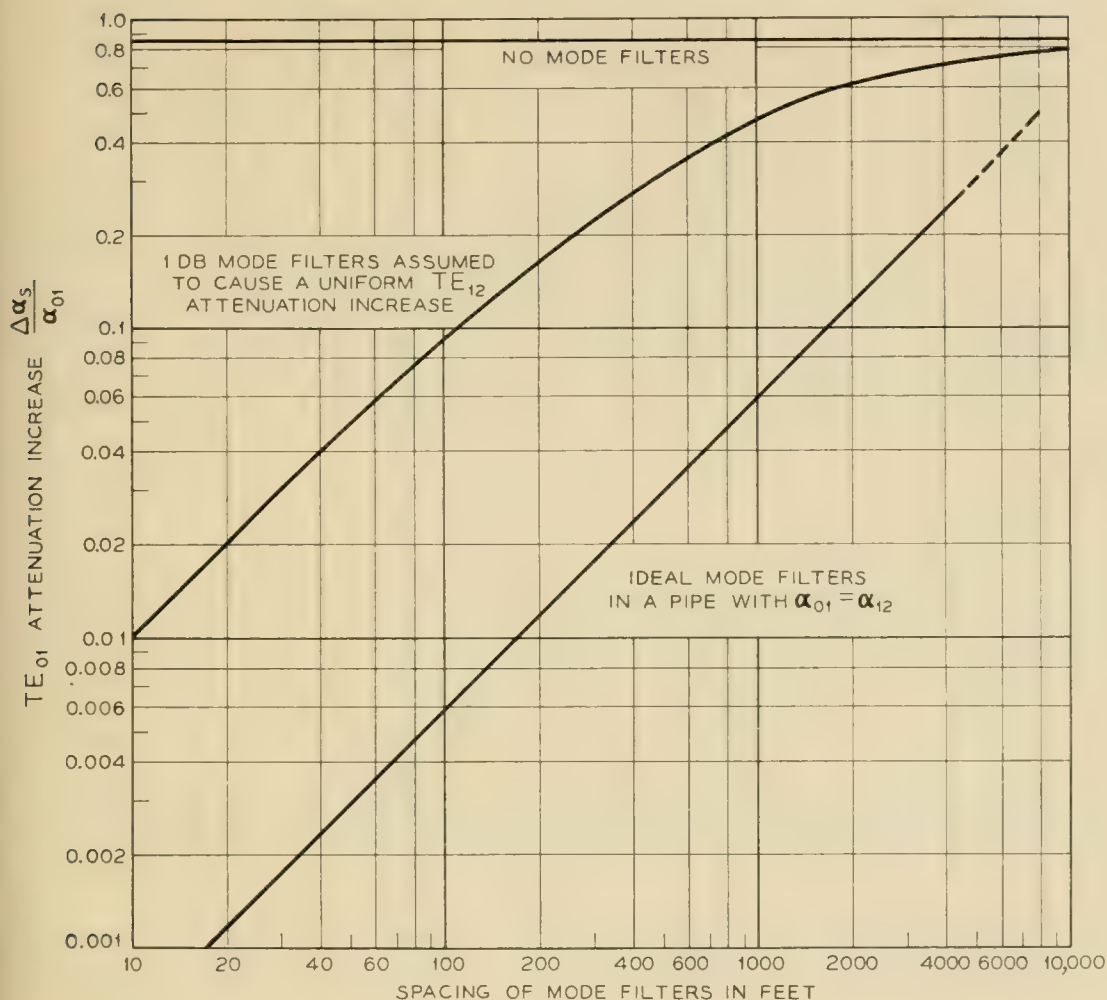


FIG. 2 — TE_{01} attenuation increase of TE_{01} - TE_{12} coupling in serpentine bends with mode filters; 2-inch I.D. 2 $\frac{3}{8}$ -inch O.D., $\lambda = 5.4$ mm. Serpentine bends are caused by equally spaced supports and the elasticity of the copper pipe. The supporting distance is a multiple of the beat-wavelength between TE_{01} and TE_{12} .

An alternative to control the coupling effects is insertion of mode filters into the line. Mode filters which pass the TE_{on} waves without loss but attenuate all other modes have been developed in various forms. To estimate the amount of mode conversion control that can be achieved by mode filters, we make two different assumptions. Only the critical case of the supporting distance equal to a multiple of the beat wavelength is considered here.

A. The mode filters are ideal; they present infinite attenuation to all unwanted modes. At the input of a section between two mode filters we have a TE_{01} wave only and at the output the spurious mode level has risen to

$$\left| \frac{E_2}{E_1} \right| = \frac{w}{EI} \frac{c_0}{(2\Delta\beta)^2} L. \quad (34)$$

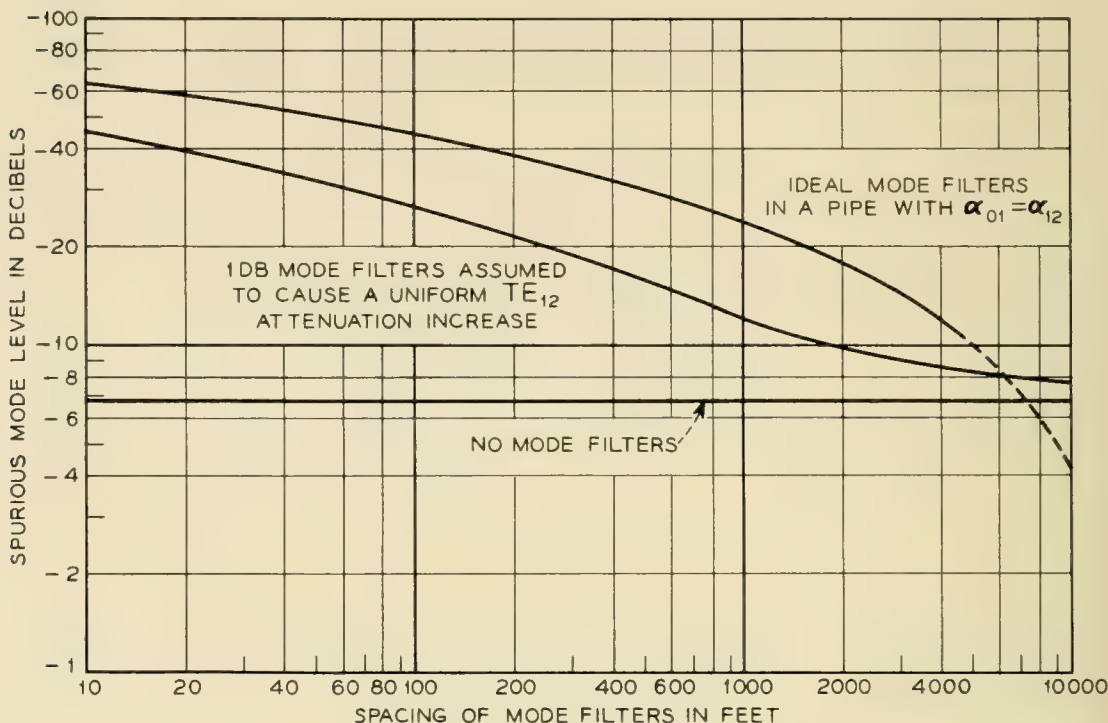


Fig. 3 — Spurious mode level of TE_{01} - TE_{12} coupling in serpentine bends with mode filters; 2-inch I.D. $2\frac{3}{8}$ -inch O.D., $\lambda = 5.4$ mm. Serpentine bends are caused by equally spaced supports and the elasticity of the copper pipe. The supporting distance is a multiple of the beat-wavelength between TE_{01} and TE_{12} .

The loss to the TE_{01} wave caused by the mode conversion is equivalent to an increase in TE_{01} attenuation

$$\Delta\alpha_s = \frac{1}{2} \left[\frac{w}{EI} \frac{c_0}{(2\Delta\beta)^2} \right]^2 L. \quad (35)$$

L is the spacing between two successive mode-filters. In (34) and (35) the attenuation constants of E_1 and E_2 are assumed to be equal. Furthermore (34) and (35) hold only as long as the E_2 power level is small compared to the E_1 power level.

Ideal mode filters is a rather optimistic assumption. Practical mode filters present only a limited attenuation to the unwanted modes. Therefore a more realistic procedure is to represent the effect of the mode filters by a uniform increase of unwanted mode attenuation.

B. The mode filters with a loss A are considered to cause a uniform increase in E_2 attenuation $\Delta\alpha = A/L$. Equations (32) and (33) modified to include this attenuation increase are:

$$\frac{\Delta\alpha_s}{\alpha_{01}} = \left[\frac{w}{EI} \frac{c_0}{(2\Delta\beta)^2 \alpha_{01}} \right]^2 \frac{\alpha_{01}}{|2\Delta\alpha| + \frac{A}{L}}, \quad (36)$$

$$\frac{E_2}{E_1} = \frac{w}{EI} \frac{c_0}{(2\Delta\beta)^2\alpha_{01}} \frac{\alpha_{01}}{|2\Delta\alpha| + \frac{A}{L}} \quad (37)$$

An evaluation of the TE_{01} - TE_{12} coupling in the previously described 2-inch copper pipe for critical frequencies near $\lambda = 5.4$ mm is shown in Figs. 2 and 3. The mode filter loss of 1 db can be achieved for the TE_{12} wave in an 18-inch long helix waveguide. A 100-foot spacing of the mode filters reduces the increase of TE_{01} attenuation to 9 per cent and the spurious mode level to -26 db.

APPENDIX

The coupling coefficient for the wave coupling between the TE_{01} wave and the TE_{11} , TE_{12} , TE_{13} waves as well as the TM_{11} wave have been calculated by S. P. Morgan.⁴ If $c = c_0/R$, the factor c_0 is for the various waves

$$TM_{11} c_0 = 0.18454\beta a,$$

$$TE_{11} c_0 = \frac{0.09319(\beta a)^2 - 0.84204}{\sqrt{\beta_{01}a\beta_{11}a}} + 0.09319 \sqrt{\beta_{01}a\beta_{11}a},$$

$$TE_{12} c_0 = \frac{0.15575(\beta a)^2 - 3.35688}{\sqrt{\beta_{01}a\beta_{12}a}} + 0.15575 \sqrt{\beta_{01}a\beta_{12}a},$$

$$TE_{13} c_0 = \frac{0.01376(\beta a)^2 - 0.60216}{\sqrt{\beta_{01}a\beta_{13}a}} + 0.01376 \sqrt{\beta_{01}a\beta_{13}a},$$

where a = radius of waveguide,

β = free-space phase constant,

β_{nm} = phase constant of TE_{nm} wave.

REFERENCES

1. W. J. Albersheim, Propagation of TE_{01} Waves in Curved Wave Guides, B.S.T.J., **28**, pp. 1-32, January, 1949.
2. S. E. Miller, Coupled Wave Theory and Waveguide Applications, B.S.T.J., **33**, pp. 661-719, May, 1954.
3. W. H. Louisell, Analysis of the Single Tapered Mode Coupler, B.S.T.J., **34**, pp. 853-870, July, 1955.
4. S. P. Morgan, Theory of Curved Circular Waveguide Containing an Inhomogeneous Dielectric, pp. 1209-1251, this issue.

Normal Mode Bends for Circular Electric Waves

By H. G. UNGER

(Manuscript received February 18, 1957)

In dielectric-coated round waveguide the degeneracy or equality of phase constants of TE_{01} and TM_{11} waves is removed. In such a non-degenerate waveguide, mode conversion in bends can be reduced by changing the curvature gradually instead of abruptly. With curvature tapers, which are of the order of, or longer than, the largest beat wavelength between TE_{01} and any of the coupled waves, propagation of only one normal mode is maintained throughout the bend. Linear curvature tapers can easily be made by bending the pipe within the limit of elastic deformation.

Changes in the direction of a waveguide line can thus be made by inserting a dielectric-coated guide section which is elastically bent over a fixed center point. A thirty degree change of direction of a 2-inch I.D. pipe with 30 ft of a dielectric-coated guide yields a total bend loss at 5.4-mm wavelength that is twice the TE_{01} loss in 30 ft of straight pipe. An optimum bend geometry is found which minimizes the total bend loss. The normal mode bend is a broadband device.

I. INTRODUCTION AND SUMMARY

A major problem in circular electric wave transmission is the question of negotiating bends. In curved sections of a round waveguide the TE_{01} mode couples to the TE_{11} , TE_{12} , TE_{13} $\cdots$ modes and to the TM_{11} mode. The coupling to the TM_{11} mode presents the most serious problem, since the TE_{01} and TM_{11} modes are degenerate in that they have equal phase velocities in a perfectly conducting straight guide. TE_{01} power introduced at the beginning of the bend will be almost completely transferred to the TM_{11} mode at odd multiples of a certain critical bending angle. This power transfer can be reduced by removing the degeneracy of equal phase velocity of TE_{01} and TM_{11} modes. There are various methods to remove the degeneracy; a simple and a very effective one is a thin dielectric layer next to the walls of the waveguide. As a study of the dielectric-

coated waveguide has shown,¹ this dielectric layer changes the phase constant of the TM_{11} wave without appreciably affecting the propagation characteristics of the TE_{01} mode.

With removal of the TE_{01} - TM_{11} degeneracy the mode conversion which occurs when a TE_{01} wave passes through a bend of constant curvature may be considerably reduced, but it will not be completely eliminated. To design a bend with still lower mode conversion losses, we consider the effect of tapering the curvature along the guide.

Guided wave propagation is most easily explained in terms of normal modes. Normal modes are solutions of the wave equation in a particular waveguide structure, and represent waves propagating without loss of power except for dissipation. In the straight waveguide the normal mode, in which we are mainly interested here is the TE_{01} mode. The normal modes of the curved section are not as simple as the straight guide modes, but they can be expressed as the sum of the normal modes in the straight waveguide. Here the mode of our main concern is the one which, when represented as a sum of straight guide modes, has most of the power in the TE_{01} part of the sum; in other words, is most similar to the TE_{01} mode of the straight waveguide.

At a transition from a straight waveguide to a curved waveguide the normal (TE_{01}) mode of the straight waveguide will certainly excite this normal mode in the curved waveguide but it will also excite a series of other normal modes. All these modes propagate in the curved section, and at the other transition to the straight guide excite not only the TE_{01} mode but a series of other normal modes of the straight waveguide. All the power in the other normal modes represents mode conversion loss of the bend.

A transition which transforms the normal (TE_{01}) mode of the straight waveguide into only one of the normal modes of the curved guide and vice versa will avoid all mode conversion losses. Such a transition can be realized approximately by tapering the curvature. Beginning with zero curvature at the straight guide end of the taper, the curvature is increased gradually along the taper to the finite value of the bend. The normal (TE_{01}) mode incident from the straight guide will be gradually transformed into the particular normal mode of the bent guide which is most similar in field configuration to the circular electric wave. At each point along the taper there is essentially only one local normal mode corresponding in its configuration to the value of curvature at that point.

This taper, unless it is infinitely long, is however only an approximation of the ideal normal mode transition. There will still be other modes excited with a very low power level. In the next section we shall analyze

the propagation in the normal mode taper. We will find that the taper should have a certain minimum length to work properly. Usually it has to be long compared to the beat wavelength between the TE_{01} normal mode and any of the other normal modes into which power may be converted.

In the plain waveguide the degeneracy between TE_{01} and TM_{11} causes an infinitely long beat wavelength. Hence, the normal mode taper would not work there. A nondegenerate waveguide is an essential condition for the normal mode bend.

We shall confine our attention to the linear taper. This is not the optimum taper form, but it is most easily built.

The residual mode conversion in the bend is to be accounted for as bend loss. This loss and the loss caused by the normal mode attenuation in the bend add up to a total bend loss. We shall evaluate the total bend loss for bend configurations which might be useful in circular electric wave transmission. For specified waveguide dimensions the total bend loss can be minimized by choosing the proper bend geometry.

The normal mode bend is an inherently broadband device. The total bend loss shows the same order of frequency dependence as the loss in the straight waveguide.

II. ANALYSIS OF THE NORMAL MODE TAPER

In the curved waveguide, wave propagation can be described in terms of the normal modes of the straight waveguide. The relation between these modes is then given by an infinite system of simultaneous first order linear differential equations. It represents the mutual coupling of the straight guide modes in the curved waveguide. We are interested mainly in TE_{01} propagation and shall restrict ourselves to a low power level in all other modes. Consequently, an adequate procedure is to consider only coupling between TE_{01} and one of the coupled modes at a time. Thus, the infinite system of equations reduces to the well known coupled line equations:²

$$\begin{aligned}\frac{dE_1}{dz} &= -\gamma_1 E_1 + jcE_2, \\ \frac{dE_2}{dz} &= jcE_1 - \gamma_2 E_2,\end{aligned}\tag{1}$$

in which

$E_{1,2}(z)$ = wave amplitudes in mode 1 (here always TE_{01}) and mode 2 (TM_{11} or one of the TE_{1m}), respectively;

$\gamma_{1,2} = j\beta_{1,2} + \alpha_{1,2}$ = propagation constants of modes 1 and 2, respectively, (the small perturbations of γ_1 and γ_2 caused by the coupling may be neglected here); and

$c(z)$ = coupling coefficient between modes 1 and 2.

In the curved waveguide the coupling coefficient is proportional to the curvature k :

$$c(z) = c'k(z) \equiv c' \frac{d\theta}{dz}, \quad (2)$$

in which θ is the direction of the guide axis. The coupled line equations (1) with varying coupling coefficient have been solved by W. H. Louisell and we shall borrow freely from his treatment.³

We define local normal modes $w_1(z)$ and $w_2(z)$:

$$\begin{aligned} E_1 &= [w_1 \cos \tfrac{1}{2} \xi - w_2 \sin \tfrac{1}{2} \xi] e^{-\gamma z}, \\ E_2 &= [w_1 \sin \tfrac{1}{2} \xi + w_2 \cos \tfrac{1}{2} \xi] e^{-\gamma z}, \end{aligned} \quad (3)$$

in which

$$\gamma = \frac{\gamma_1 + \gamma_2}{2},$$

$$\tan \xi = j2 \frac{c}{\gamma_2 - \gamma_1} \equiv j2 \frac{c}{\Delta\gamma}.$$

Substituting (3) into (1), we find that $w_1(z)$ and $w_2(z)$ must satisfy

$$\begin{aligned} \frac{dw_1}{dz} - \Gamma(z)w_1 &= \frac{1}{2} \frac{d\xi}{dz} w_2, \\ \frac{dw_2}{dz} + \Gamma(z)w_2 &= -\frac{1}{2} \frac{d\xi}{dz} w_1, \end{aligned} \quad (4)$$

where $\Gamma(z) = \frac{1}{2} \sqrt{\Delta\gamma^2 - 4c^2}$. In (4) the local normal modes are coupled only through the terms proportional to $d\xi/dz$. When ξ is constant they are uncoupled and true normal modes. For small values of $d\xi/dz$ or more specifically when

$$\left| \frac{1}{2\Gamma} \frac{d\xi}{dz} \right| \ll 1 \quad (5)$$

approximate solutions of (4) can be written down, which proceed essentially in powers of $d\xi/dz$. These solutions are:

$$\begin{aligned}
 w_1(z) = e^{\rho(z)} & \left[w_1(0) + \frac{1}{2} w_2(0) \int_0^z \frac{d\xi}{dz'} e^{-2\rho(z')} dz' \right. \\
 & \left. - \frac{1}{4} w_1(0) \int_0^z \frac{d\xi}{dz'} e^{-2\rho(z')} \int_0^{z'} \frac{d\xi}{dz''} e^{2\rho(z'')} dz'' dz' \right], \\
 w_2(z) = e^{-\rho(z)} & \left[w_2(0) - \frac{1}{2} w_1(0) \int_0^z \frac{d\xi}{dz'} e^{2\rho(z')} dz' \right. \\
 & \left. - \frac{1}{4} w_2(0) \int_0^z \frac{d\xi}{dz'} e^{2\rho(z')} \int_0^{z'} \frac{d\xi}{dz''} e^{-2\rho(z'')} dz'' dz' \right],
 \end{aligned} \tag{6}$$

in which

$$\rho(z) = \int_0^z \Gamma(z') dz'.$$

The initial conditions in the normal mode taper are $E_1(0) = 1$ and $E_2(0) = 0$. The taper begins with zero curvature, $\xi(0) = 0$. Hence from (3) $w_1(0) = 1$ and $w_2(0) = 0$. The wanted local normal mode is w_1 while w_2 is an unwanted mode. At the end, z_1 , of the taper the unwanted mode

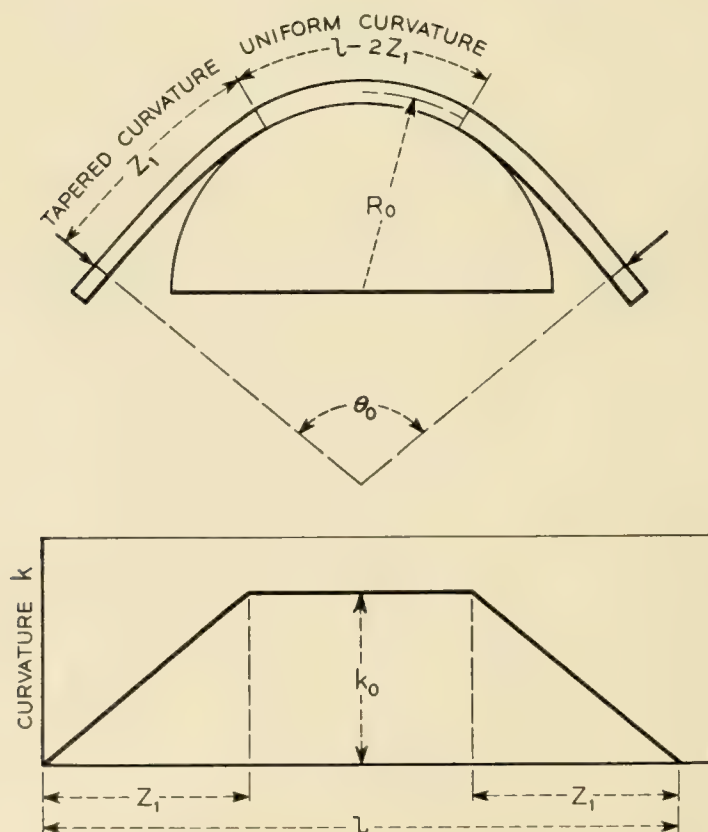


Fig. 1 — Normal mode bend with linear curvature taper.

amplitude is

$$|w_2(z_1)| = \frac{1}{2} \left| \int_0^{z_1} \frac{d\xi}{dz} e^{2\rho(z)} dz \right|. \quad (7)$$

This amplitude represents mode conversion loss and therefore has to be kept as small as possible.

In (7) the function $\xi(z)$, i.e., the taper function, is still undetermined. Obviously it can be chosen so as to optimize the taper performance. A taper of optimal design keeps the unwanted mode below a certain value with as short a taper length as possible. From (7) the relation between this optimizing problem and the problem of the transmission line taper of optimal design⁴ is at once evident. The transmission line taper is a low reflection transition between lines of different characteristic impedances. To minimize the length of the transition for a specified maximal reflection, the characteristic impedance has to change along the transition according to a function which is essentially the Fourier transform of a Tschebyscheff polynomial of infinite degree. The same procedure can be applied here and it will result in a curvature taper of optimal design.

We are, however, at present not as much interested in a transition of optimal design as in a curvature taper which can easily be built. Suppose we bend the pipe to a bending radius R_o which causes only elastic deformation. We do this on a form of radius R_o , Fig. 1. The forces acting on both ends of the pipe cause a torque and hence a curvature of the pipe which increases approximately linearly from the pipe end ($z = 0$) to the point of contact ($z = z_1$) between pipe and form:

$$k = k_0 \frac{z}{z_1}. \quad (8)$$

The corresponding curve which the pipe forms along the taper is Cornu's spiral.

We shall evaluate (7) for a curvature as given by (8). In considering the mode conversion we may neglect all heat losses, that is $\gamma = j\beta$, etc. With

$$c = c_0 \frac{z}{z_1}, \quad (9)$$

we get

$$\frac{d\xi}{dz} = \frac{2c_0}{\Delta\beta z_1} \frac{1}{1 + 4 \frac{c_0^2 z^2}{\Delta\beta^2 z_1^2}}, \quad (10)$$

and

$$\rho(z) = j \frac{\Delta\beta^2 z_1}{4c_0} \left[\frac{c_0 z}{\Delta\beta z_1} \sqrt{1 + 4 \left(\frac{c_0 z}{\Delta\beta z_1} \right)^2} + \frac{1}{2} \sinh^{-1} \frac{2c_0 z}{\Delta\beta z_1} \right]. \quad (11)$$

We introduce (10) and (11) into (7) and take advantage of

$$2 \left| \frac{c_0}{\Delta\beta} \right| \ll 1, \quad (12)$$

which is satisfied for a gentle enough bend. The unwanted mode level is then given by

$$|w_2(z_1)| = \left| \frac{c_0}{\Delta\beta^2 z_1} (1 - e^{j\Delta\beta z_1}) \right| + 0 \left(8 \frac{c_0^3}{\Delta\beta^3} \right). \quad (13)$$

The general restriction (5) for the solution (6) in case of the linear taper is

$$\frac{2c_0}{\Delta\beta^2 z_1} \left(1 + 4 \frac{c_0^2 z^2}{\Delta\beta^2 z_1^2} \right)^{-3/2} \ll 1,$$

and in view of (12) only

$$|\Delta\beta z_1| \geq 1 \quad (14)$$

is required. The length of the transition has to be of the order or greater than the largest beat wavelength.

In addition to mode conversion loss the normal mode suffers heat loss along the taper. From (3) and (6) this heat loss is given by the real part of $[\gamma z - \rho(z)]$. If $\Delta\beta \gg \Delta\alpha$ we get

$$R[\gamma z - \rho] = \alpha_1 z + \Delta\alpha \int_0^z \frac{c^2}{\Delta\beta^2} dz'. \quad (15)$$

It follows from (15) that the attenuation of the coupled straight guide modes should not be too high for the normal mode bend to work properly.

III. THE TOTAL BEND LOSS

We will consider bends of the dielectric-coated waveguide only. The normal modes of the dielectric-coated waveguide have phase constants which are slightly different from the phase constants of the modes in the plain guide¹

$$\begin{aligned} \text{TM}_{nm} \quad \frac{\Delta\beta}{\beta_{nm}} &= \frac{\epsilon' - 1}{\epsilon'} \delta, \\ \text{TE}_{nm} \quad \frac{\Delta\beta}{\beta_{nm}} &= \frac{n^2}{p_{nm}^2 - n^2} \frac{\epsilon' - 1}{\epsilon'(1 - \nu^2)} \delta, \\ \text{TE}_{0m} \quad \frac{\Delta\beta}{\beta_{0m}} &= \frac{p_{0m}^2}{3} \frac{\epsilon' - 1}{1 - \nu_{0m}^2} \delta^3. \end{aligned} \quad (16)$$

The losses in the dielectric coat increase the attenuation constants by:

$$\begin{aligned} \text{TM}_{nm} \frac{\alpha_D}{\beta_{nm}} &= \frac{\epsilon''}{\epsilon'^2} \delta, \\ \text{TE}_{nm} \frac{\alpha_D}{\beta_{nm}} &= \frac{n^2}{p_{nm}^2 - n^2} \frac{\epsilon''}{\epsilon'^2 (1 - \nu_{nm}^2)} \delta, \\ \text{TE}_{0m} \frac{\alpha_D}{\beta_{0m}} &= \frac{p_{0m}^2}{3} \frac{\epsilon''}{1 - \nu_{0m}^2} \delta^3. \end{aligned} \quad (17)$$

For the circular electric waves the dielectric coat increases the wall current attenuation by

$$TE_{0m} \frac{\Delta\alpha}{\alpha_{0m}} = (\epsilon' - 1) \frac{p_{0m}^2}{\nu_{0m}^2} \delta^2. \quad (18)$$

The various symbols in (16), (17) and (18) are:

β_{nm} = plain guide phase constants of TM_{nm} and TE_{nm} respectively;
 p_{nm} = m th root of $J_n(x) = 0$ for TM_{nm} , and m th root of $J_n'(x) = 0$ for TE_{nm} waves respectively;

$\nu_{nm} = \frac{\lambda}{\lambda_c} = \frac{p_{nm}\lambda}{2\pi a}$ cutoff factor in plain waveguide of radius a ;

$\delta = \frac{d}{a}$ relative thickness of dielectric coat;

$\epsilon = \epsilon' - j\epsilon''$ relative dielectric permittivity of dielectric coat;

α_{0m} = attenuation constant of TE_{0m} in the plain waveguide.

The coupling coefficient between the straight guide modes in a curved waveguide is $c = c'/R$ in which:⁵

$$\text{TE}_{01} \leftrightarrow \text{TM}_{11} c' = 0.18454 \beta a,$$

$$\text{TE}_{01} \leftrightarrow \text{TE}_{11} c' = \frac{0.09319(\beta a)^2 - 0.84204}{\sqrt{\beta_{01}a\beta_{11}a}} + 0.09319 \sqrt{\beta_{01}a\beta_{11}a},$$

$$\text{TE}_{01} \leftrightarrow \text{TE}_{12} c' = \frac{0.15575(\beta a)^2 - 3.35688}{\sqrt{\beta_{01}a\beta_{12}a}} + 0.15575 \sqrt{\beta_{01}a\beta_{12}a}, \quad (19)$$

$$\text{TE}_{01} \leftrightarrow \text{TE}_{13} c' = \frac{0.01376(\beta a)^2 - 0.60216}{\sqrt{\beta_{01}a\beta_{13}a}} + 0.01376 \sqrt{\beta_{01}a\beta_{13}a},$$

where β = free-space phase constant.

We consider the bend configuration of Fig. 1, with the curvature being a trapezoidal function of length. The maximum power loss due to conversion to one of the unwanted modes in the first transition is, by (13),

$$|w_2(z_1)|_{\max}^2 = \frac{4}{(\Delta\beta_{z_1})^2} \frac{c_0^2}{\Delta\beta^2}.$$

By the law of reciprocity the same conversion loss occurs in the second transition. Both conversion parts phase with each other. To get an average total conversion we may add them in quadrature and the total conversion loss to one of the unwanted modes is, expressed in nepers,

$$A_c = \frac{4}{(\Delta\beta z_1)^2} \frac{c_0^2}{\Delta\beta^2}. \quad (20)$$

Under most unfavorable phase conditions the conversion loss may be twice this value, but it is very unlikely that such phase conditions will be satisfied for all coupled modes simultaneously.

Besides mode conversion loss the local normal mode suffers attenuation in the bend. This attenuation is larger than the straight guide attenuation. Each straight guide mode contained in the local normal mode of the bend causes an increase in attenuation. From (15) the loss caused by one of these straight guide modes is:

$$A_b = (\alpha_{1m} - \alpha_{01}) \int_0^l \frac{c^2}{\Delta\beta^2} dz. \quad (21)$$

Where α_{1m} is the attenuation constant of the TE_{1m} and TM_{11} waves respectively in the dielectric coated waveguide.

Introducing the trapezoidal curvature function of Fig. 1 into (21) we get

$$A_b = \frac{c_0^2}{\Delta\beta^2} (\alpha_{1m} - \alpha_{01}) \left(l - \frac{4}{3} z_1 \right). \quad (22)$$

The loss caused by the TE_{01} attenuation in the straight dielectric coated guide is from (17) and (18)

$$A_s = \alpha_{01} l \left[1 + (\epsilon - 1) \frac{p_{0m}^2}{\nu_{0m}^2} \delta^2 \right] + \frac{p_{01}^2}{3} \frac{\epsilon'' \beta_{01}}{1 - \nu_{01}^2} \delta^3 l. \quad (23)$$

The total bend loss is finally obtained by summing up all the terms of (20), (22), and (23),

$$A = A_s + \sum A_b + \sum A_c. \quad (24)$$

The summation signs indicate that all coupled modes (TM_{11} and TE_{1m}) have to be taken into account.

The effectiveness of the normal mode bend is best demonstrated by a practical example. A copper pipe now in experimental use at Bell Telephone Laboratories for circular electric wave transmission near 5.4 mm wavelength has 2-inch I. D. and $2\frac{3}{8}$ -inch O.D. Suppose we want to change the direction of a waveguide line with this copper pipe by an angle θ_0 . We can do this most easily by inserting a dielectric-coated section, which

is bent around a fixed support in the center by forces acting on both ends. In order not to exceed elastic deformation the bending radius must not be smaller than

$$R_{\min} = \frac{E}{f_{\max}} a_1, \quad (25)$$

where $f_{\max}$ = flexural stress at elastic limit,
 E = modulus of elasticity,
 a_1 = outside radius of pipe.

This minimum bending radius requires a minimum length to change the pipe direction by a specified angle θ_0 given by

$$l_{\min} = 2\theta_0 R_{\min}. \quad (26)$$

The total bend loss (24) has been evaluated for a bend configuration as specified by (25) and (26). The result is shown in Fig. 2. The total additional bend loss is only of the order of the TE_{01} loss in the plain straight waveguide. For small bending angles the curvature taper becomes shorter and consequently the mode conversion loss increases. The mode conversion loss, however, does not go to infinity for zero bending angle. In this case (14) is no longer satisfied, and the mode conversion loss is no longer described by (20).

The level of the various unwanted modes which can be calculated from (20) is plotted in Fig. 2.

For a practical waveguide one would decide on a standard length of dielectric-coated pipe, one or several of which would be inserted whenever a change in direction has to be made. Take, in our example, a standard length of 15 feet. With one such section a change of direction up to 15° could be made. For a change in direction up to 30° two such sections would have to be inserted and bent around a fixed support at the center joint. The total loss of Fig. 2 is then a maximum value, which would only occur when the pipe is bent to the highest allowable bending angle.

IV. A NORMAL MODE BEND OF OPTIMUM DESIGN

The various terms of the total bend loss (24) depend on the bend geometry in quite different ways. It is therefore likely that for a given bending angle θ_0 a bend geometry can be found, which minimizes the total bend loss. The total bend loss can generally be written as:

$$A = Sl + B \frac{1 - \frac{4}{3}u}{l(1 - u)^2} + C \frac{1}{l^4(u - u^2)^2}, \quad (27)$$

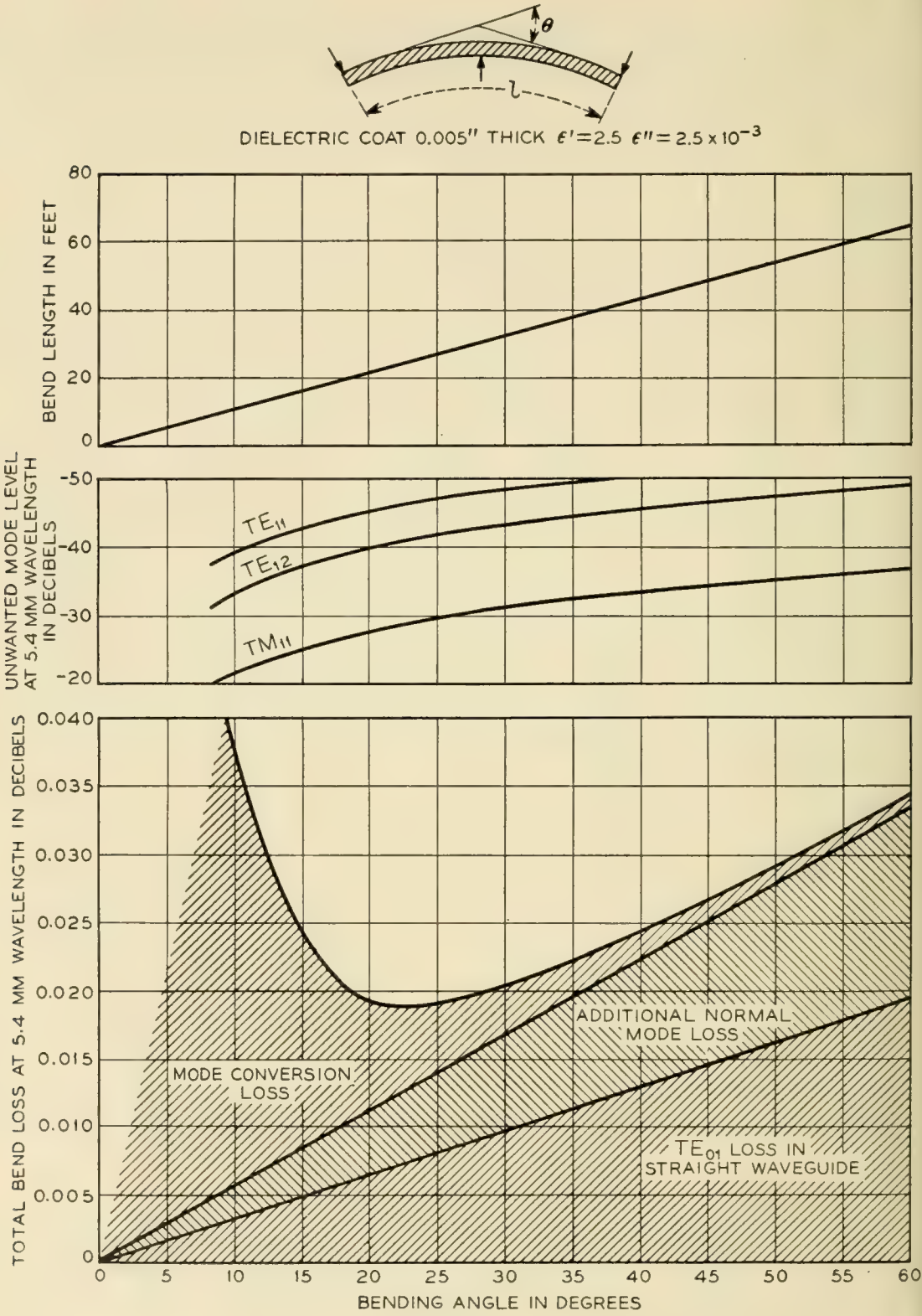


Fig. 2 — Normal mode bend. Dielectric-coated copper pipe with 2-inch I.D. and 2 $\frac{3}{8}$ -inch O.D. deflected to limit of elastic deformation ($\lambda = 5.4$ mm).

in which

$$\begin{aligned} S &= \alpha_{01} \left[1 + (\epsilon' - 1) \frac{p_{01}^2}{\nu_{01}^2} \delta^2 \right] + \frac{p_{01}^2}{3} \frac{\epsilon'' \beta_{01}}{1 - \nu_{01}^2} \delta^3, \\ B &= \Sigma \frac{c'^2}{\Delta \beta^2} \theta_0^2 (\alpha_{1m} - \alpha_{01}), \\ C &= \Sigma \frac{c'^2}{\Delta \beta^4} \theta_0^2, \\ u &= \frac{z_1}{l}. \end{aligned} \tag{28}$$

Here again the summation signs indicate that all coupled modes have to be taken into account. The factors S , B , and C do not depend on the bend geometry, but only on the total bending angle, the waveguide properties, and the frequency. Necessary conditions for $A(u, l)$ to be a minimum are:

$$\frac{\partial A}{\partial u} \equiv \frac{2(1 - 2u)}{l(1 - u^3)} \left[\frac{B}{3} - \frac{C}{l^3 u^3} \right] = 0, \tag{29}$$

with the two roots

$$u = \frac{1}{2}, \tag{30}$$

$$ul = \left(3 \frac{C}{B} \right)^{\frac{1}{3}}, \tag{31}$$

and

$$\frac{\partial A}{\partial l} \equiv S - B \frac{1 - \frac{4}{3}u}{(1 - u)^2} \frac{1}{l^2} - \frac{4C}{(u - u^2)^2} \frac{1}{l^5} = 0. \tag{32}$$

If $u = \frac{1}{2}$, the solutions of (32) are the roots of

$$S l^5 - \frac{4}{3} B l^3 - 64C = 0. \tag{33}$$

If $(lu)^3 = 3(C/B)$, the solutions of (32) are the roots of

$$l = \left(3 \frac{C}{B} \right)^{\frac{1}{3}} \pm \left(\frac{B}{S} \right)^{\frac{1}{2}}. \tag{34}$$

Sufficient conditions for $A(u, l)$ to be a minimum are:

$$\frac{\partial^2 A}{\partial u^2} \frac{\partial^2 A}{\partial l^2} - \left(\frac{\partial^2 A}{\partial u \partial l} \right)^2 > 0, \tag{35}$$

$$\frac{\partial^2 A}{\partial u^2} > 0, \quad \frac{\partial^2 A}{\partial l^2} > 0, \tag{36}$$

If $u = \frac{1}{2}$, we have:

$$\begin{aligned}\frac{\partial^2 A}{\partial u \partial l} &\equiv 0, \\ \frac{\partial^2 A}{\partial l^2} &= \frac{8}{l^3} \left(\frac{B}{3} + 40 \frac{C}{l^3} \right), \\ \frac{\partial^2 A}{\partial u^2} &= \frac{32}{l} \left(8 \frac{C}{l^3} - \frac{B}{3} \right).\end{aligned}\tag{37}$$

Substituting

$$\begin{aligned}x &= \frac{1}{2} \left(\frac{S}{B} \right)^{\frac{1}{2}} l, \\ r &= \left(3 \frac{C}{B} \right)^{\frac{1}{2}} \left(\frac{S}{B} \right)^{\frac{1}{2}},\end{aligned}\tag{38}$$

we get instead of (33)

$$x^3(3x^2 - 1) - 2r^3 = 0,\tag{39}$$

and instead of (37)

$$\frac{\partial^2 A}{\partial u^2} = 16 \frac{B}{l} (x^2 - 1).$$

The positive root of (39) is plotted in Fig. 3. It follows that if $r > 1$ we have $x > 1$, consequently $\partial^2 A / \partial u^2 > 0$; and if $r < 1$ we have $x < 1$, consequently $\partial^2 A / \partial u^2 < 0$. Consequently if, and only if, $r > 1$ the values $u = \frac{1}{2}$ and x from Fig. 3 minimize the total bend loss A.

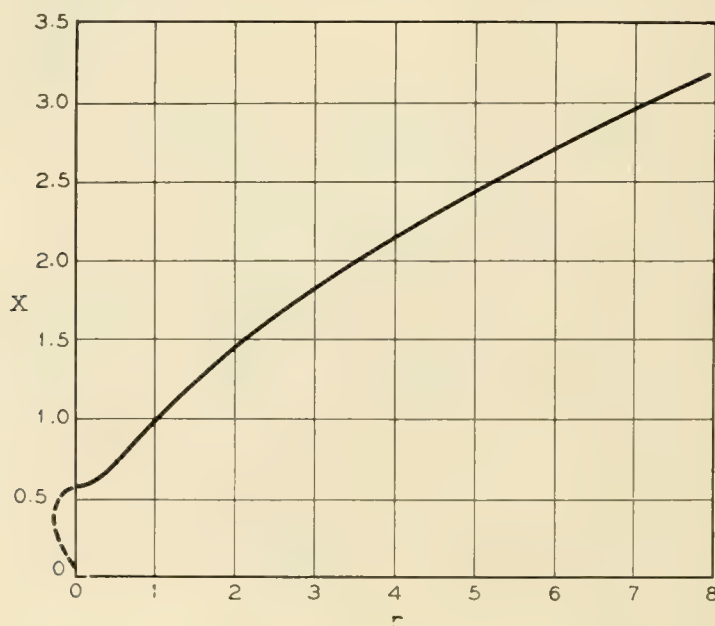


Fig. 3 — Positive root of $x^3(3x^2 - 1) - 2r^3 = 0$.

If $(ul)^3$ is equal to $3(C/B)$,

$$\frac{\partial^2 A}{\partial u^2} \frac{\partial^2 A}{\partial l^2} - \left(\frac{\partial^2 A}{\partial u \partial l} \right)^2 \equiv \frac{4B^2}{l^4 u} \frac{1 - 2u}{(1 - u)^6},$$

and

$$\frac{\partial^2 A}{\partial u^2} \equiv \frac{2B}{lu} \frac{1 - 2u}{(1 - u)^3}.$$

Hence, if $u < \frac{1}{2}$ or, because of (31) and (34) $r < 1$, a minimum of $A(u, l)$ is located at

$$(ul)^3 = 3 \frac{C}{B}, \quad l = \left(3 \frac{C}{B} \right)^{\frac{1}{3}} + \left(\frac{B}{S} \right)^{\frac{1}{3}}.$$

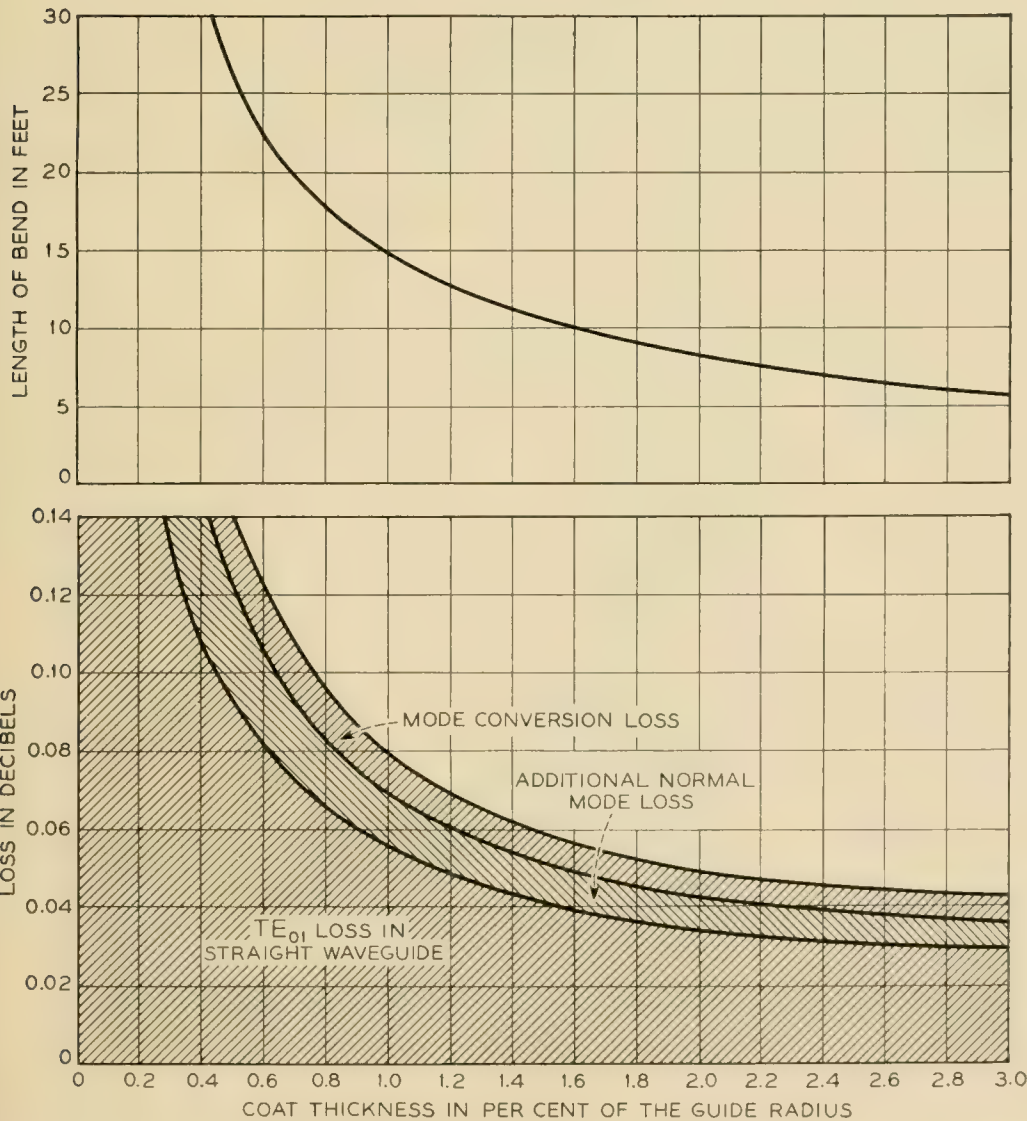


Fig. 4 — 90° normal mode bend of $\frac{7}{8}$ -inch I.D. copper pipe with a dielectric coat of $\epsilon' = 2.5$, $\epsilon'' = 2.5 \times 10^{-3}$. Optimum design for 5.4-mm wavelength.

To find the optimum bend geometry for a given dielectric coated guide and a specified bending angle we calculate r from (38). If $r > 1$ the optimum geometry is

$$l = 2 \sqrt{\frac{B}{S}} x \quad \text{and} \quad z_1 = \sqrt{\frac{B}{S}} x$$

with x from Fig. 3. If $r < 1$,

$$l = \left(3 \frac{C}{B} \right)^{\frac{1}{3}} + \left(\frac{B}{S} \right)^{\frac{1}{2}}, \quad \text{and} \quad z_1 = \left(3 \frac{C}{B} \right)^{\frac{1}{3}}.$$

A numerical example, the 90° bend of a $\frac{7}{8}$ -inch I. D. copper pipe, is shown in Fig. 4. The total bend loss in the optimally designed bend decreases steadily with increasing thickness of the dielectric coat. This indicates that there is also an optimum coat thickness, which minimizes the total loss of the normal mode bend of optimum total length and taper length. Unfortunately, however, several approximations made in calculating phase constants and coupling coefficients in the dielectric-coated

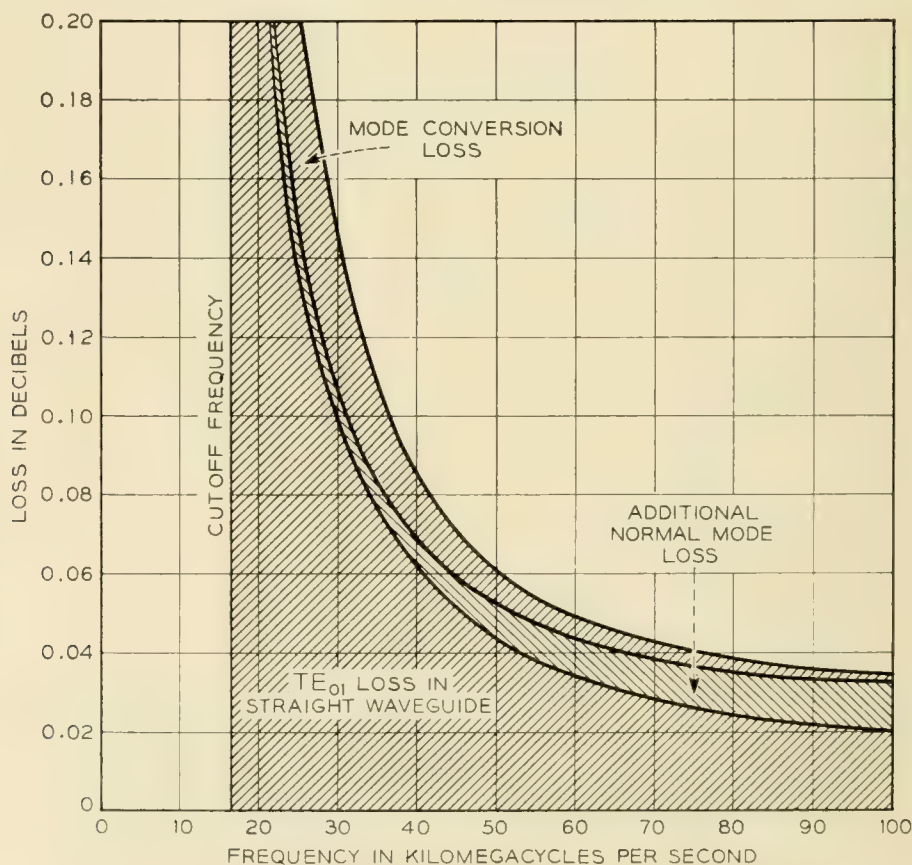


Fig. 5 — 90° normal mode bend of $\frac{7}{8}$ -inch I.D. copper pipe with a dielectric coat 0.0075 inch thick ($\epsilon' = 2.5$, $\epsilon'' = 2.5 \times 10^{-3}$), designed for optimum performance at $\lambda = 5.4$ mm (55.5 kmc).

waveguide usually break down at smaller than optimum values of the coat thickness.

It should be mentioned finally that the normal mode bend is an inherently broad band device. Except for the oscillations of the mode conversion portion of the total loss as caused by spurious mode phasing, there is only a gradual change of the loss with frequency.

Some terms contributing to the total loss decrease with frequency, others increase. The over-all frequency dependence is of the same order as the frequency dependence of the loss in the straight waveguide. As an example, in Fig. 5 the bend loss has been plotted versus frequency for the normal mode bend of Fig. 4.

ACKNOWLEDGMENT

Mathematical analysis of tapered curvature in other forms of waveguide has been made by others. S. E. Miller reports that Siemens & Halske A. G., Germany, have made an original treatment of this subject.

REFERENCES

1. H. G. Unger, Circular Electric Wave Transmission in the Dielectric Coated Waveguide, pp. 1253-1278, this issue.
2. S. E. Miller, Coupled Wave Theory and Waveguide Applications, B.S.T.J., **33**, pp. 661-719, May, 1954.
3. W. H. Louisell, Analysis of the Single Tapered Mode Coupler, B.S.T.J., **34**, pp. 853-870, July, 1955.
4. R. W. Klopfenstein, A Transmission Line Taper of Improved Design, Proc. I.R.E., **44**, pp. 31-35, January, 1956.
5. S. P. Morgan, Theory of Curved Circular Waveguide Containing an Inhomogeneous Dielectric, pp. 1209-1251, this issue.

Bell System Technical Papers Not Published in This Journal

ASHKIN, A.¹

Electron Beam Analyzer, J. Appl. Phys., **28**, p. 564, May, 1957.

BECK, A. C.¹

Communications Superhighways, Trans. I.R.E., PGM TT, MTT-5, pp. 81-82, April, 1957.

BECKER, G. E.¹

Dependence of Magnetron Operation on the Radial Centering of the Cathode, Trans. I.R.E., PGED, ED-4, pp. 126-131, April, 1957.

BOOTHBY, O. L., see Williams, H. J.

BOWERS, F. K.¹

What Use is Delta Modulation to the Transmission Engineer? Commun. & Electronics, **30**, pp. 142-147, May, 1957.

BROWN, S. C., see Rose, D. J.

BROYER, A. P., see Schlabach, T. D.

BUCK, T. M.,¹ and McKIM, F. S.¹

Experiments on the Photomagnetolectric Effect in Germanium. Phys. Rev., **106**, pp. 904-909, June 1, 1957.

BUEHLER, E.¹

Contribution to the Floating Zone Refining of Silicon, Rev. Sci. Instr., **28**, pp. 453-460, June, 1957.

CHILBERG, G. L.²

Buried Cable Telephone Systems, Commun. & Electronics, **30**, pp. 130-135, May, 1957.

¹ Bell Telephone Laboratories, Inc.

² American Telephone and Telegraph Company.

CHYNOWETH, A. G.,¹ and MCKAY, K. G.¹

Internal Field Emission in Silicon P-N Junctions, *Phys. Rev.*, **106**, pp. 418-426, May 1, 1956.

CLEMENCY, W. F.,¹ ROMANOW, F. F.,¹ and ROSE, A. F.²

The Bell System Speakerphone, *Commun. & Electronics*, **30**, pp. 148-153, May, 1957.

COMPTON, K. G.¹

Variability in Working Copper Sulfate Half Cells, *Corrosion*, **13**, pp. 19-20, March, 1957.

COWAN, F. A.²

Transmission of Color Over Nationwide Television Networks, *S.M.P.T.E.*, **66**, pp. 278-283, May, 1957.

DRENICK, R.¹

An Operational View of Equipment Reliability, *Trans. Annual Convention, Am. Soc. Quality Control*, **11**, pp. 603-611, May 22, 1957.

EASLEY, J. W.¹

The Effect of Collector Capacity on the Transient Response of Junction Transistors, *Trans. I.R.E., PGED, ED-4*, pp. 6-14, Jan., 1957.

EMLING, J. W.¹

General Aspects of Hands-Free Telephony, *Commun. & Electronics*, **30**, pp. 201-205, May, 1957.

FAGEN, R. E., see Hall, A. D.

FELDMANN, W. L., see Pearson, G. L.

FORD, B. W.⁵

Demand Increasing for Special Uses of Telephone Facilities, *Telephony*, **152**, pp. 22-24, 44, 48, June 8, 1957.

FREERICKS, L.¹

Semiconductors in Switching Circuits, *General Engineering Bulletin (N. Y. Tel. Co.)*, **7**, pp. 21-25, May, 1957.

¹ Bell Telephone Laboratories, Inc.

² American Telephone and Telegraph Company.

⁵ Illinois Bell Telephone Company, Chicago.

FREUDENSTEIN, F.,¹ WARTHMAN, K. L.,¹ and WATROUS, A. B.¹

Designing Gear-Train Limit Stops for Control of Shaft Rotation, Machine Design, **11**, pp. 84-86, May 30, 1957.

GALT, J. K.¹

Losses in Ferrites — Single Crystal Studies, J. Inst. Elec. Engrs. (London), **104**, pp. 189-197, June, 1957.

GELLER, S.¹

Comments on Pauling's Paper on Effective Metallic Radii for Use in the β -Wolfram Structure, Acta Cryst., **10**, pp. 380-382, May 10, 1957.

GERARD, H. B.¹

Some Effects of Status, Role Clarity and Group Goal Clarity Upon the Individual's Relations to Group Process, J. Personality, **25**, pp. 475-488, June, 1957.

GREEN, E. I.¹

Evaluating Scientific Personnel, Elec. Engg., **76**, pp. 578-584, July, 1957.

GUPTA, S. S.,¹ HUYETT, M. J.,¹ and SOBEL, M.¹

Selection and Ranking Problems with Binomial Populations, Trans. Annual Convention, Am. Soc. Quality Control, **11**, pp. 635-718, May 22, 1957.

HAKE, E. A.¹

A 10-Kw Germanium Rectifier for Automatic Power Plants, A.I.E.E. Conf. Publication "Rectifiers in Industry", **T-93**, p. 119, June, 1957.

HALL, A. D.,¹ and FAGEN, R. E.¹

Definition of System, Yearbook Soc. Advancement of General Systems Theory, **1**, pp. 18-28, April, 1957.

HARKER, K. J.¹

Non-Laminar Flow in Cylindrical Electron Beams, J. Appl. Phys., **28**, pp. 645-650, June, 1957.

HENNESSEY, T. M.⁶

A Dial Service and Community Relationships. Telephony, **152**, pp. 22-24, 40, June 1, 1957.

¹ Bell Telephone Laboratories, Inc.

⁶ New England Telephone and Telegraph Company, Boston, Mass.

HERRMANN, G.¹

Transverse Scaling of Electron Beams, J. Appl. Phys., **28**, pp. 474-478, April, 1957.

HUYETT, M. J., see Gupta, S. S.

LEVENBACH, G. J.¹

Accelerated Life Testing of Capacitors, Trans. I.R.E., PGRQC, RQC-10, pp. 9-20, June, 1957.

LLOYD, S. M.³

A New Water Rheostat for Testing Exchange Power Plants, Telephony, **152**, pp. 32-34, May 25, 1957.

LUMSDEN, G. Q.¹

Wood Poles for Communication Lines, A.S.T.M. Bulletin, **222**, pp. 19-24, May, 1957.

McCALL, D. W.¹

Cell for the Determination of Pressure Coefficients of Dielectric Constant and Loss of Liquids and Solids to 10,000 psi, Rev. Sci. Instr., **28**, pp. 345-351, May, 1957.

McCALL, D. W., see Slichter, W. P.

McKAY, K. G., see Chynoweth, A. G.

McKIM, F. S., see Buck, T. M.

McSKIMIN, H. J.¹

Use of High Frequency Ultrasound for Determining the Elastic Moduli of Small Specimens, Proc. National Electronics Conf., **12**, pp. 351-362, April 15, 1957.

MILLER, W. A.⁷

A Narrow-Band Experimental FM Mobiletelephone System, Commun. & Electronics, **30**, pp. 98-100, May, 1957.

MONTGOMERY, H. C.¹

Field Effect in Germanium at High Frequencies, Phys. Rev., **106**, pp. 441-445, May 1, 1957.

¹ Bell Telephone Laboratories, Inc.

³ Western Electric Company, Inc.

⁷ Pacific Telephone and Telegraph Company, Los Angeles, Calif.

O'BRIEN, J. A.¹

Unit Distance Binary-Decimal Code Translators, Trans. I.R.E., PGEC, EC-6, pp. 122-123, June, 1957, Letter to the Editor.

PEARSON, G. L.,¹ READ, W. T., JR.,¹ and FELDMANN, W. L.¹

Deformation and Fracture of Small Silicon Crystals, Acta Met., 5, pp. 181-191, April, 1957.

PEDERSON, C. W.¹

Crystal Clock for Airborne Computer, Electronics, 30, pp. 196-198 June 1, 1957.

READ, W. T., JR., see Pearson, G. L.

RIDER, D. K., see Schlabach, T. D.

ROMANOW, F. F., see Clemency, W. F.

ROSE, A. F., see Clemency, W. F.

ROSE, D. J.,¹ and BROWN, S. C.⁴

Microwave Gas Discharge Breakdown in Air, Nitrogen, and Oxygen, J. Appl. Phys., 28, pp. 561-563, May, 1957.

SCHLABACH, T. D.,¹ WRIGHT, E. E.,¹ BROYER, A. P.,¹ and RIDER, D. K.¹

Testing of Foil-Clad Laminates for Printed Circuitry, Bull A.S.T.M., 222, pp. 25-30, May, 1957.

SHENITZER, A.¹

On the Problem of Chebyshev Approximation of a Continuous Function by a Class of Functions, J. Assoc. Computing Machinery, 4, pp. 30-35, Jan., 1957.

SHERWOOD, R. C., see Williams, H. J.

SNOKE, L. R.¹

Some Needed Basic Research on Wood Deterioration Problems, Appl. Microbiology, 5, pp. 188-193, May, 1957.

SOBEL, M., see Gupta, S. S.

¹ Bell Telephone Laboratories, Inc.

⁴ Massachusetts Institute of Technology, Cambridge.

VAN UITERT, L. G.¹

Effects of Annealing on the Saturation Induction of Ferrites Containing Nickel and/or Copper, J. Appl. Phys., **28**, pp. 478-481, April, 1957.

WARTHMAN, K. L., see Freudenstein, F.

WATROUS, A. B., see Freudenstein, F.

WHITTEMORE, L. E.²

The Institute of Radio Engineers — Forty-five Years of Service, Proc. I.R.E., **45**, pp. 597-635, May, 1957.

WILLIAMS, H. J.,¹ and SHERWOOD, R. C.¹

Magnetic Domain Patterns on Thin Films, J. Appl. Phys., **28**, pp. 548-555, May, 1957.

WILLIAMS, H. J.,¹ SHERWOOD, R. C.¹ and BOOTHBY, O. L.¹

Magnetostriction and Magnetic Anisotropy of MnBi, J. Appl. Phys., **28**, pp. 445-447, April, 1957.

WRIGHT, E. E., see Schlabach, T. D.

¹ Bell Telephone Laboratories, Inc.

² American Telephone and Telegraph Company.

Recent Monographs of Bell System Technical Papers Not Published in This Journal*

BENNETT, W. R.

Synthesis of Active Networks, Monograph 2816.

DEWALD, J. F.

Formation of Anode Films on Single-Crystal Indium Antimonide, Monograph 2802.

DITZENBERGER, J. A., see Fuller, C. S.

EIGLER, J. H., see Sullivan, M. V.

* Copies of these monographs may be obtained on request to the Publication Department, Bell Telephone Laboratories, Inc., 463 West Street, New York 14, N. Y. The numbers of the monographs should be given in all requests.

FLASCHEN, S. S., see Garn, P. D.

FULLER, C. S., and DITZENBERGER, J. A.

Effect of Defects in Germanium on Diffusion and Acceptance of Copper, Monograph 2803.

FULLER, C. S., and MORIN, F. J.

Diffusion and Electrical Behavior of Zinc in Silicon, Monograph 2789.

GARN, P. D., and FLASCHEN, S. S.

Analytical Applications of a Differential Thermal Analysis Apparatus, Monograph 2769.

GEBALLE, T. H., see Kunzler, J. E.

GEBALLE, T. H., see Morin, F. J.

GOHN, G. R., see Torrey, M. N.

GOULD, H. L. B., and WENNY, D. H.

Supremendur: A New Rectangular-Loop Magnetic Material, Monograph 2780.

GREEN, E. I.

Nature's Pulses, Monograph 2770.

GREEN, E. I.

Science and Liberal Education, Monograph 2765.

GREEN, E. I.

Telephone, Monograph 2806.

HAGSTRUM, H. D.

Auger Ejection of Electrons from Tungsten by Noble Gas Ions, Monograph 2715.

HAGSTRUM, H. D.

Effect of Monolayer Adsorption on Ejection of Electrons from Metals, Monograph 2804.

HAGSTRUM, H. D.

Thermionic Constants and Sorption Properties of Hafnium, Monograph 2790.

HERRING, C., see Morin, F. J.

HULL, G. W., see Kunzler, J. E.

KLEMM, G. H.

Automatic Protection Switching for TD-2 Radio System, Monograph 2772.

KOHN, W.

Effective Mass Theory in Solids from a Many-Particle Standpoint, Monograph 2791.

KUNZLER, J. E., GEBALLE, T. H., and HULL, G. W.

Germanium Resistance Thermometers Suitable for Low-Temperature Calorimetry, Monograph 2792.

LLOYD, S. P., and McMILLAN, B.

Linear Least Squares Filtering and Prediction of Sampled Signals, Monograph 2815.

LUNDBERG, C. V., see Vacca, G. N.

McMILLAN, B., see Lloyd, S. P.

MENDIZZA, A.

Standard Salt-Spray Test — Is It A Valid Acceptance Test?, Monograph 2807.

MESZAR, J.

Full Stature of the Crossbar Tandem Switching System, Monograph 2750.

MILLER, S. L.

Ionization Rates for Holes and Electrons in Silicon, Monograph 2794.

MORIN, F. J., see Fuller, C. S.

MORIN, F. J., GEBALLE, T. H., and HERRING, C.

Temperature Dependence of Piezoresistance of High-Purity Silicon, Germanium, Monograph 2797.

MORIN, F. J., and REISS, H.

Formation of Ion Pairs and Triplets Between Lithium and Zinc in Germanium, Monograph 2795.

REISS, H., see Morin, F. J.

ROSE, D. J.

Microplasmas in Silicon, Monograph 2796.

SMITH, K. D., see Veloric, H. S.

SUHL, H.

Theory of Ferromagnetic Resonance at High Signal Powers, Monograph 2817.

SULLIVAN, M. V., and EIGLER, J. H.

Electroless Nickel Plating for Making Ohmic Contacts to Silicon, Monograph 2808.

SULLIVAN, M. V., and EIGLER, J. H.

Five Metal Hydrides as Alloying Agents on Silicon, Monograph 2775.

TORREY, M. N., and GOHN, G. R.

A Study of Statistical Treatments of Fatigue Data, Monograph 2778.

VACCA, G. N., and LUNDBERG, C. V.

Aging of Neoprene in a Weatherometer, Monograph 2809.

VAN HORN, R. H.

Experimental Evaluation of Reliability of Solderless Wrapped Connections, Monograph 2810.

VELORIC, H. S., and SMITH, K. D.

Silicon Diffused Junction "Avalanche" Diodes, Monograph 2811.

VOGEL, F. L., JR.

Dislocations in Plastically Bent Germanium Crystals, Monograph 2763.

WEBER, L. A.

Influence of Noise on Telephone Signaling Circuit Performance, Monograph 2812.

WEIBEL, E. S.

An Electronic Analog Multiplier Using Carriers, Monograph 2813.

WENNY, D. H., see Gould, H. L. B.

YOUNKER, E. L.

A Transistor-Driven Magnetic-Core Memory, Monograph 2814.

Contributors to This Issue

C. H. ELMENDORF, B.S., California Institute of Technology, 1935; M.S., 1936; Bell Telephone Laboratories, 1936-. Mr. Elmendorf was concerned with the development of coaxial cable transmission systems from 1936 to 1955, except for the period from 1941-1945 when he worked on airborne radar systems. In 1952 he started an exploratory program on submarine cable systems. As Assistant Director of Transmission Systems Development since 1955, he has been responsible for development of submarine cable systems. He is a senior member of the I.R.E.

BRUCE C. HEEZEN, B.A., University of Iowa, 1948; M.A., Columbia University, 1952; Ph.D. 1956. After participating in a cruise of the Woods Hole research vessel *Atlantis* in the summer of 1948, Mr. Heezen held a fellowship in geology at Columbia, where he joined the staff of the Lamont Geological Observatory when it was founded in 1949. As submarine geologist, and now as senior scientist in charge of the submarine geology program, he has been a member of numerous deep-sea expeditions. In addition, he teaches a graduate course in submarine geology on the Columbia campus. His work includes deep-sea topography, sediments, and sedimentation processes; submarine photography and deep-sea research instrumentation; and geologic, geophysical, and oceanographic exploration of the deep sea.

SAMUEL P. MORGAN, B.S., 1943; M.S. and Ph.D., 1947, California Institute of Technology; Bell Telephone Laboratories, 1947-. A research mathematician, Mr. Morgan has been concerned with the application of electromagnetic theory to microwave problems, and has also made studies in other fields of mathematical physics. Member American Physical Society, Tau Beta Pi, Sigma Xi and I.R.E.

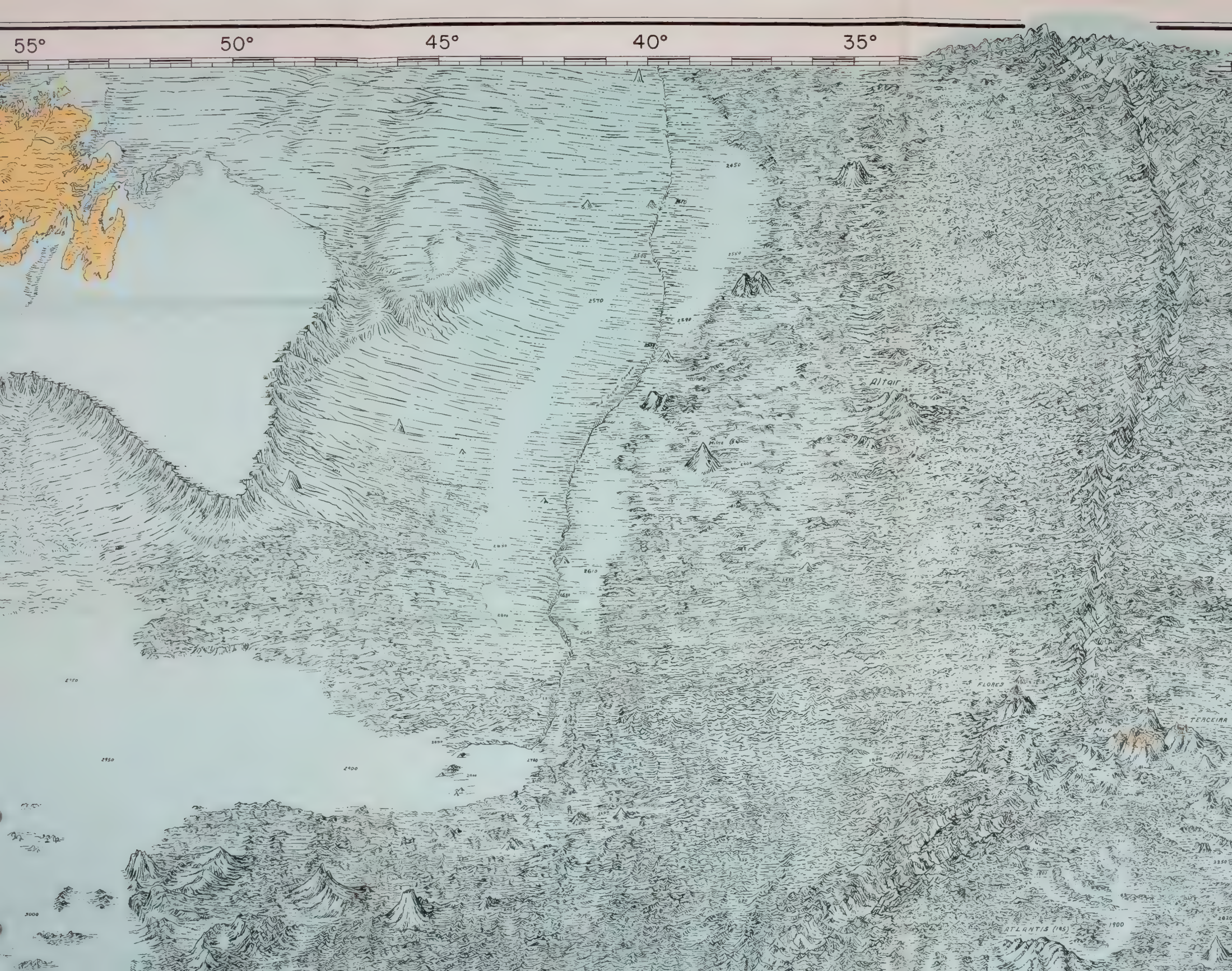
LLOYD R. SNOKE, B.S. in For., 1948, Pennsylvania State University; Bell Telephone Laboratories, 1948-. Since joining the Laboratories, Mr. Snoke has specialized in the timber products used in the Bell System and their preservative treatment. He has been specifically engaged in the study of timber treatment theory, the application of radioactive isotopes to fundamental problems and the bioassay of wood preservatives. For the past four years Mr. Snoke has been concerned with microbiological testing of materials including laboratory bacteriological studies and actual marine tests. He heads the Environmental Protection Group of the Outside Plant Development Department. Member American Association for the Advancement of Science, the Society for Industrial Microbiology, American Wood Preservers' Association, American Society for Testing Materials, Materials Advisory Board — Technical Panel on Miscellaneous Materials, Steering Committee of Microbiological Deterioration Section — Gordon Research Conferences and Zi Sigma Pi.

HANS-GEORG UNGER, Dipl. Ing., 1951, Dr. Ing., 1954, Technische Hochschule Braunschweig (Germany); Bell Telephone Laboratories 1956-. Mr. Unger's work at the Laboratories has been in waveguides, especially circular electric wave transmission. He holds several foreign patents on waveguides and has published in German technical magazines. Member I.R.E.

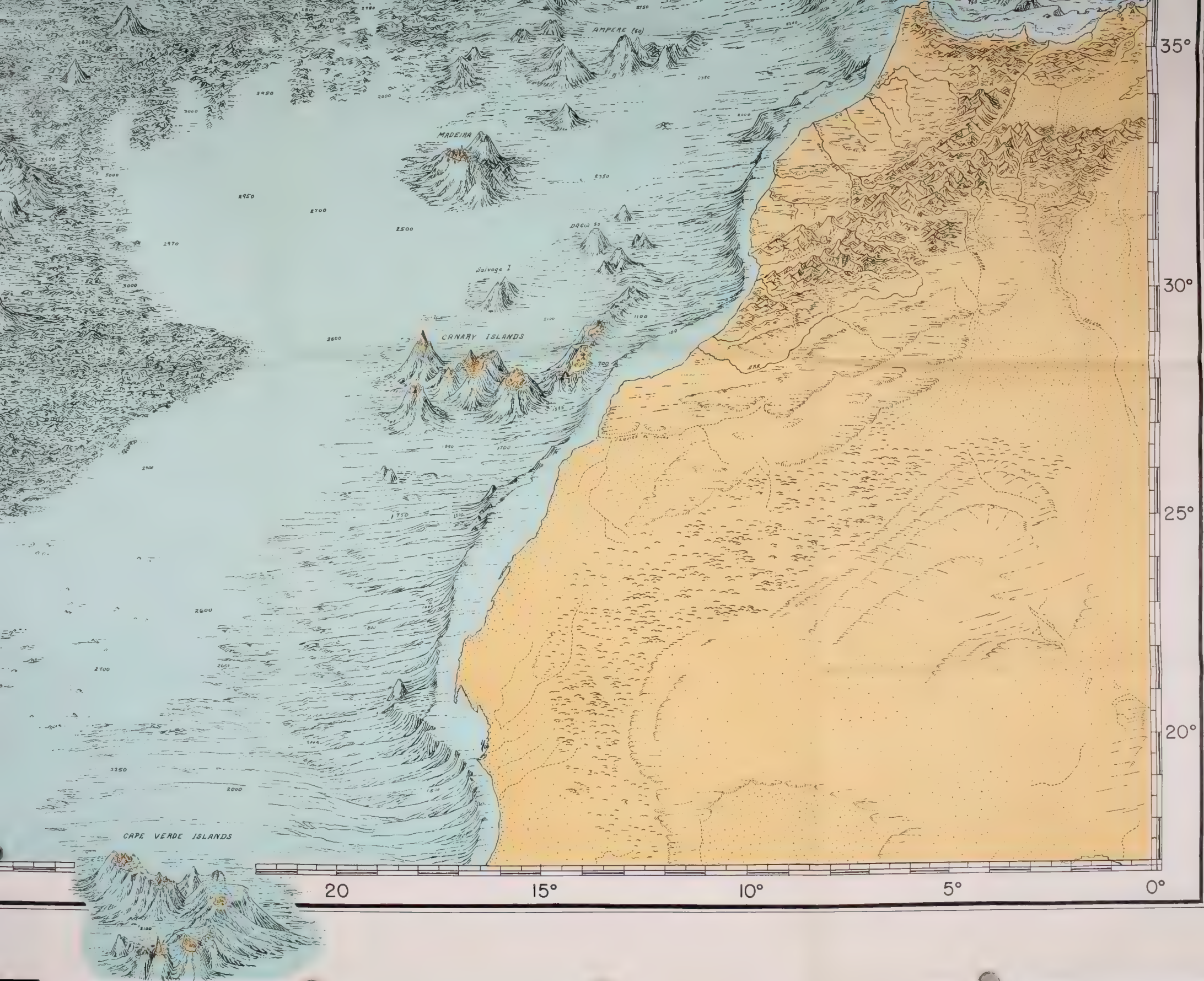
E. E. ZAJAC, B.M.E., 1950, Cornell University; M.S.E., 1952, Princeton University; Ph.D., 1954, Stanford University; Bell Telephone Laboratories, 1954-. Since joining the Laboratories, Mr. Zajac's work has been in theoretical and applied mechanics in the Mathematical Research Department. Member of the American Society of Mechanical Engineers, Tau Beta Pi, Pi Tau Sigma, Phi Kappa Phi and Sigma Xi.

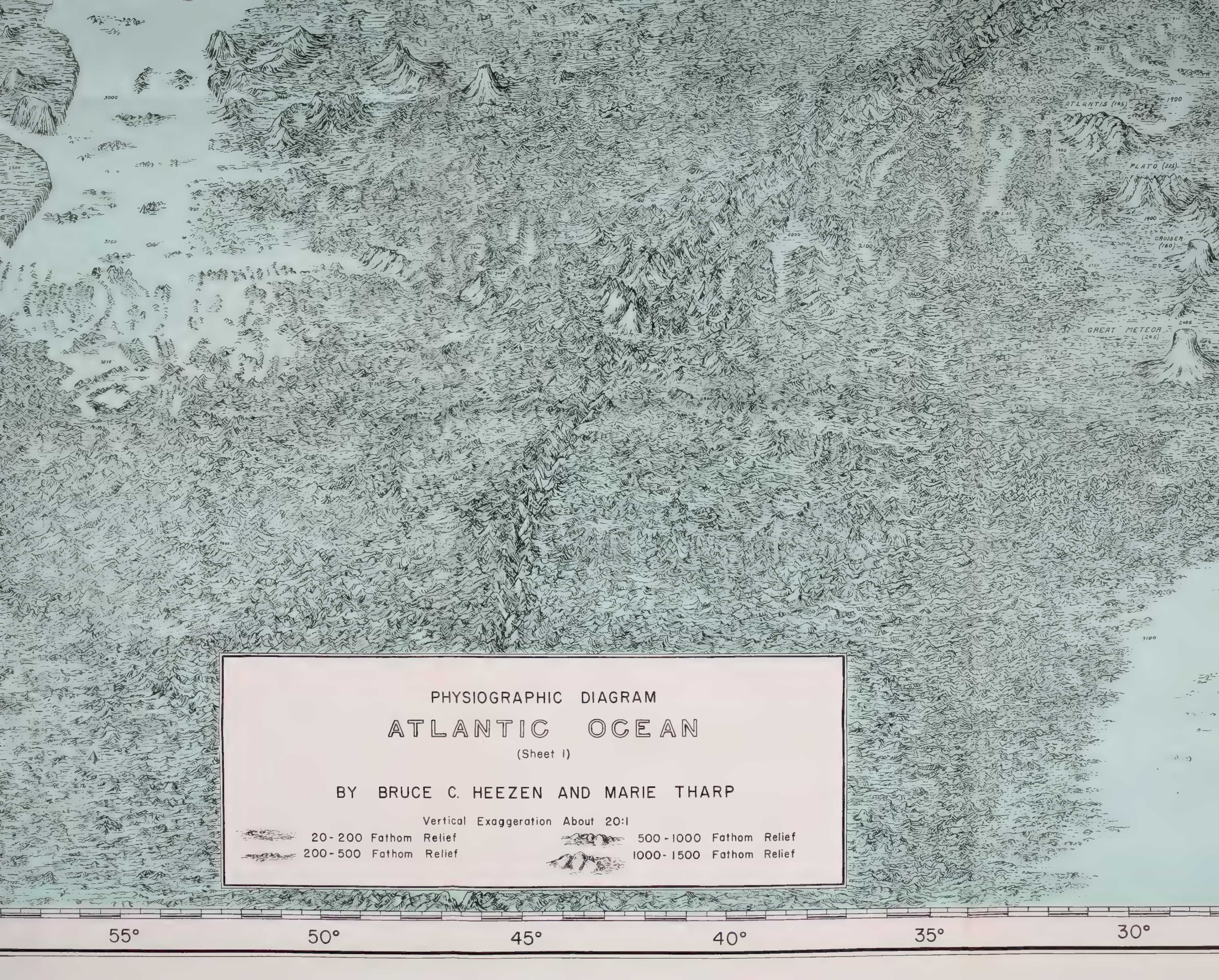
PHYSIOGRAPHIC DIAGRAM OF THE ATLANTIC OCEAN
SEE ELMENDORF AND HEEZEN ARTICLE
PAGES 1047-1093







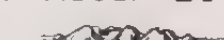




PHYSIOGRAPHIC DIAGRAM
ATLANTIC OCEAN
(Sheet I)

BY BRUCE C. HEEZEN AND MARIE THARP

Vertical Exaggeration About 20:1

	20-200 Fathom Relief		500-1000 Fathom Relief
	200-500 Fathom Relief		1000-1500 Fathom Relief

55°

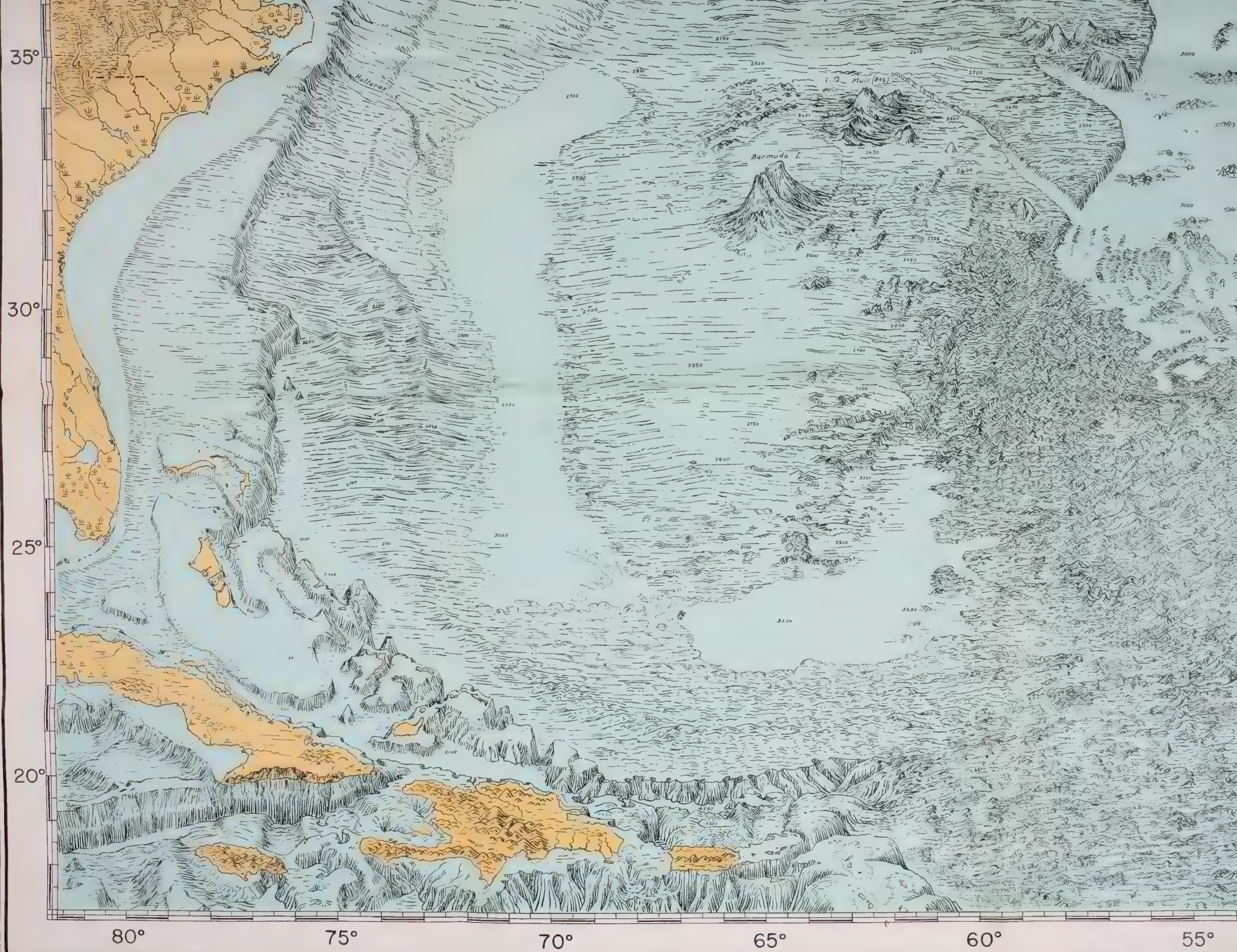
50°

45°

40°

35°

30°



B4T

THE BELL SYSTEM

Technical Journal

DEVOTED TO THE SCIENTIFIC AND ENGINEERING
ASPECTS OF ELECTRICAL COMMUNICATION

VOLUME XXXVI

NOVEMBER 1957

NUMBER 6

A New Storage Element Suitable for Large-Sized Memory Arrays
—The Twistor A. H. BOBECK 1319

Non-Binary Error Correction Codes WERNER ULRICH 1341

Shortest Connection Networks and Some Generalizations
R. C. PRIM 1389

A Network Containing a Periodically Operated Switch Solved by
Successive Approximations C. A. DESOER 1403

Experimental Transversal Equalizer for TD-2 Radio Relay
System B. C. BELLOWS AND R. S. GRAHAM 1429

Transmission Aspects of Data Transmission Service Using Private
Line Voice Telephone Channels P. MERTZ AND D. MITCHELL 1451

Design, Performance and Application of the Vernier Resolver
G. KRONACHER 1487

Bell System Technical Papers Not Published in This Journal 1501

Recent Bell System Monographs 1508

Contributors to This Issue 1511

THE BELL SYSTEM TECHNICAL JOURNAL

ADVISORY BOARD

A. B. GOETZE, *President, Western Electric Company*

M. J. KELLY, *President, Bell Telephone Laboratories*

E. J. MCNEELY, *Executive Vice President, American Telephone and Telegraph Company*

EDITORIAL COMMITTEE

B. McMILLAN, *Chairman*

K. E. GOULD

S. E. BRILLHART

E. I. GREEN

A. J. BUSCH

R. K. HONAMAN

L. R. COOK

H. R. HUNTLEY

A. C. DICKINSON

F. R. LACK

R. L. DIETZOLD

J. R. PIERCE

EDITORIAL STAFF

W. D. BULLOCH, *Editor*

R. L. SHEPHERD, *Production Editor*

T. N. POPE, *Circulation Manager*

THE BELL SYSTEM TECHNICAL JOURNAL is published six times a year by the American Telephone and Telegraph Company, 195 Broadway, New York 7, N. Y. F. R. Kappel, President; S. Whitney Landon, Secretary; John J. Scanlon, Treasurer. Subscriptions are accepted at \$5.00 per year. Single copies \$1.25 each. Foreign postage is 65 cents per year or 11 cents per copy. Printed in U. S. A.

THE BELL SYSTEM TECHNICAL JOURNAL

VOLUME XXXVI

NOVEMBER 1957

NUMBER 6

Copyright 1957, American Telephone and Telegraph Company

A New Storage Element Suitable for Large-Sized Memory Arrays— The Twistor

By ANDREW H. BOBECK

Three methods have been developed for storing information in a coincident-current manner on magnetic wire. The resulting memory cells have been collectively named the "twistor". Two of these methods utilize the strain sensitivity of magnetic materials and are related to the century old Wertheim or Wiedemann effects; the third utilizes the favorable geometry of a wire.

The effect of an applied torsion on a magnetic wire is to shift the preferred direction of magnetization into a helical path inclined at an angle of 45° with respect to the axis. The coincidence of a circular and a longitudinal magnetic field inserts information into this wire in the form of a polarized helical magnetization. In addition, the magnetic wire itself may be used as a sensing means with a resultant favorable increase in available signal since the lines of flux wrap the magnetic wire many times. Equations concerning the switching performance of a twistor are derived.

An experimental transistor-driven, 320-bit twistor array has been built. The possibility of applying weaving techniques to future arrays makes the twistor approach appear economically attractive.

I. INTRODUCTION

A century ago Wiedemann¹ observed that if a suitable magnetic rod which carries a current is magnetized by an external axial field, a twist

of the rod will result. The effect is a consequence of the resultant helical flux field causing a change in length of the rod in a helical sense. Conversely, it was also observed that a rod under torsion will produce a voltage between its ends when the rod is magnetized (see Fig. 1).

Recently, during an investigation of the magnetic properties of nickel wire, it was observed that a voltage was developed across the ends of a nickel wire as its magnetization state was changed. Both the amplitude and the polarity of the observed signal could be varied by movements of the nickel wire. Most surprising, the amplitude of the observed voltage v_2 of Fig. 2, was many times that which would be expected if a conventional pickup loop were used.

After determining experimentally that the observed voltage was generated solely in the nickel wire and was not a result of air flux coupling the sensing loop (nickel wire plus unavoidable copper return wire), it was concluded that the flux in the nickel wire must follow a helical path. This suggested that torsion was the cause of the observed effect, a conclusion verified experimentally. The direction of the applied twist de-

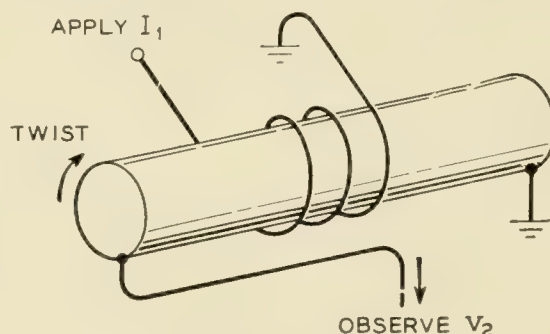


Fig. 1 — Observation of an internally induced voltage v_2 generated by a magnetic wire under torsion.

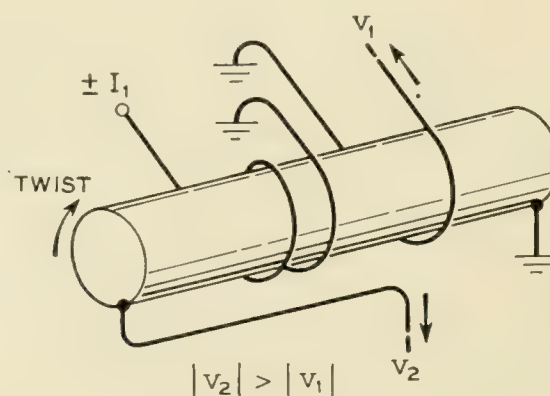


Fig. 2 — Comparison of the internally induced voltage v_2 to the voltage v_1 induced in the pickup loop.

terminated the polarity and the amount of the twist determined the magnitude of the observed voltage.

As a consequence of these results, it is possible to build mechanical-to-electrical transducers,² transformers with unity turns ratio but possessing a substantial transforming action, and a variety of basic memory cells.

This paper will be concerned with a discussion of the memory cells from both a practical and theoretical viewpoint. It will be shown how these cells can be fabricated into memory arrays. One such configuration consists solely of vertical copper wires and horizontal magnetic wires. Experimental results of the switching behavior of many magnetic materials when operated in the "twisted" manner will be given.

II. A COINCIDENT-CURRENT MEMORY CELL — THE TWISTOR

Consider a wire rigidly held at the far end and subjected to a clockwise torsion applied to the near end. This will result in a stress component of maximum compression³ at an angle of 45° with respect to the axis of the wire in the right-hand screw sense, and a component of maximum tension following a left-hand screw sense. All magnetic materials are strain sensitive to some degree. This will depend upon both the chemical composition and the mechanical working of the material. For example, if unannealed nickel wire is subjected to a torsion, the preferred direction of magnetization will follow the direction of greatest compression, as would be predicted from the negative magnetostrictive coefficient of nickel. Unannealed nickel wire, then, will have a preferred remanent flux path as shown in Fig. 3.

If the ease of magnetization as measured along the helix is sufficiently lower than that along the axis or circumference, it is possible to insert information into the wire in a manner somewhat analogous to the usual

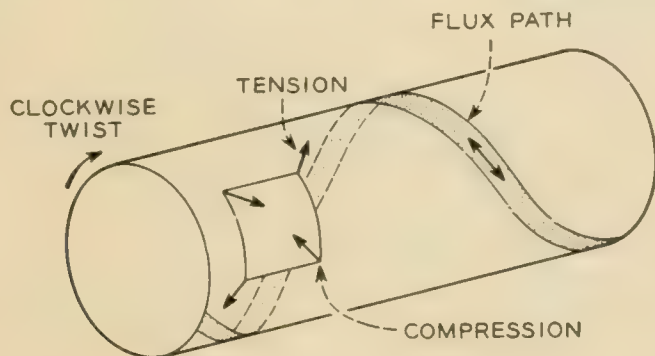


Fig. 3 — Relationship of the mechanical stresses resulting from applied torsion to the preferred magnetic flux path in nickel.

coincident-current method. Consider a current pulse I_1 applied through the nickel wire in such a direction as to enhance the spiraling flux, and a second current pulse I_2 applied by means of an external solenoid (see Fig. 4). Coincidentally, the proper amplitude current pulses will switch the flux state of the wire; either alone will not be sufficient. To sense the state of the stored information it is necessary either to reverse both currents, or to overdrive I_2 in the reverse direction. In an array, the output, in the form of a voltage pulse, would be sensed across the ends of the nickel wire. The solenoid may be replaced by a single copper conductor passing at right angles to the nickel wire. For obvious reasons the memory cell has been named the "twistor". The above method of operation will be referred to as mode A.

Mode B is the use of the magnetic wire as a direct replacement for the conventional coincident-current toroid. Its use here differs only in that the wire itself is used as a sensing winding (refer to Fig. 5). The pulses I_1 and I_2 are equal in value and each alone is chosen to be insufficient to switch the magnetization state of the wire. The coincidence of I_1 and I_2 will, however, result in the writing of a bit of information into the wire. To read, I_1 and I_2 are reversed in polarity and applied coincidentally. The output appears as a voltage pulse across the ends of the nickel wire.

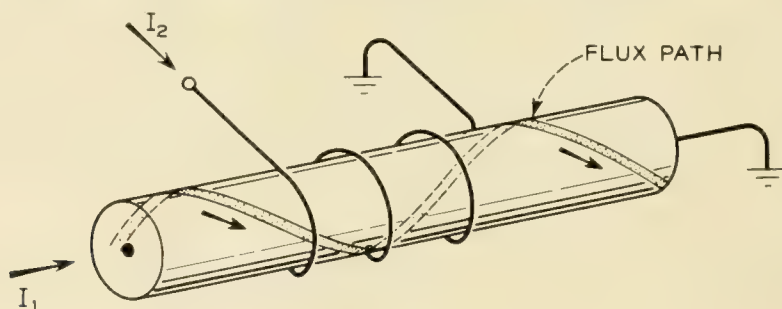


Fig. 4 — Coincident currents for the "write" operation in a twistor operated mode A. Wire under torsion.

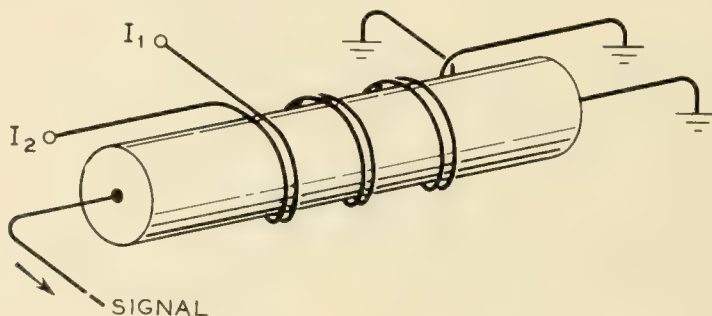


Fig. 5 — For mode B the coincidence of I_1 and I_2 is required to exceed the knee of the φ - NI characteristic. Wire under torsion.

The third method of operating the twistor, mode C, is similar in nature to a method proposed by J. A. Baldwin.⁴ In this scheme the wire is not twisted, so that neither screw sense is favorable. By the proper application of external current pulses, information will be stored in the wire in the form of a flux path of a right-hand screw sense for a "1", and a left-hand screw sense for a "0". The operation of the cell is indicated in Figure 6. Note that the writing procedure requires a coincidence of currents; the reading procedure does not. The sign of the output voltage indicates the stored information.

Modes A and C are best suited for moderate sized memory arrays since the reading procedure is not a coincident type selection. Thus to gain access to n^2 storage points, an access switch capable of selecting one of n^2 points is required. For large arrays the use of mode B is indicated. It then becomes possible to select one of n^2 points with a $2n$ position access switch. The crossover point (about 10^5 bits) is determined by access circuitry considerations.

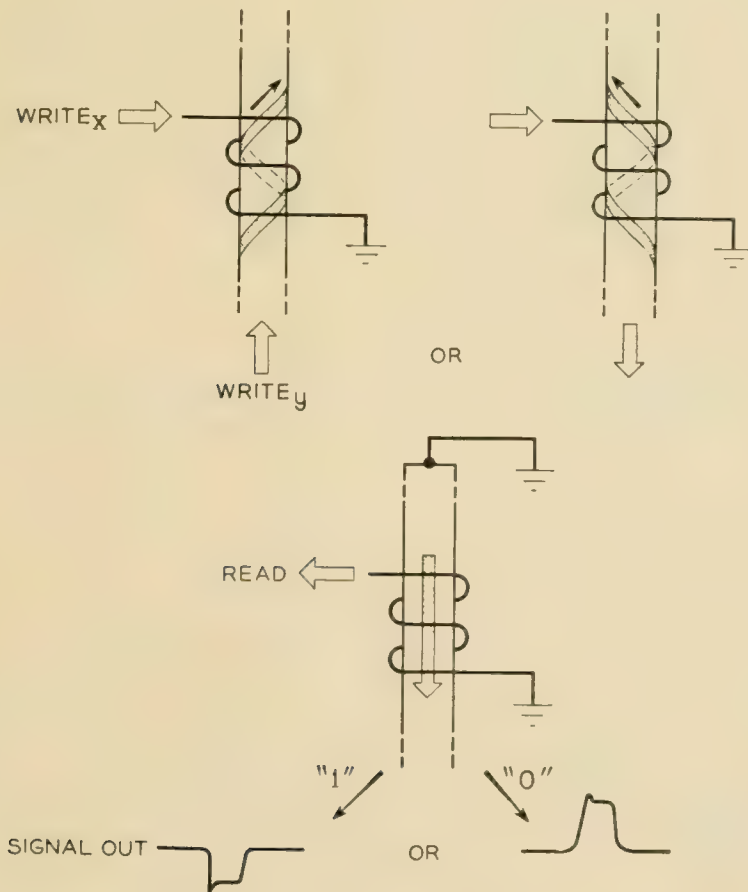


Fig. 6 — Read-write cycle for a twistor operated mode C. The wire is not under torsion.

III. ANALYSIS OF THE SWITCHING PROPERTIES OF THE TWISTOR

Section 3.1 will deal with the basic properties of magnetic wire as they pertain to the twistor memory cell. Section 3.2 will be concerned with a composite magnetic wire. The theoretical conclusions will be supported by experimental results wherever possible.

3.1 *Solid Magnetic Wire*

It has been stated above that there is a voltage gain inherent in the operation of the twistor. This voltage gain makes it possible to obtain millivolt signals from wires several mils in diameter. An expression will now be derived relating the axial flux of an *untwisted* wire to the circular flux component of that wire when twisted. Assume that the magnetic wire has been twisted so that the flux spirals at an angle θ (normally $\theta = 45^\circ$) with respect to the axis of the wire. If d and l are the diameter and length of the magnetized region respectively, then, for a complete flux reversal the change in the *circular* flux component is $\varphi_{\text{circ}} = l(d/2)(2B_s \sin \theta)$.

Here, φ_{circ} is the flux change that would be observed on a hypothetical pickup wire which passed down the axis of the magnetic wire. The flux change which would be observed by a single pickup loop around the wire, if the magnetic wire were not twisted, is $\varphi_{\text{longitudinal}} = \pi d^2 B_s / 2$. Therefore, $\varphi_{\text{circ}} / \varphi_{\text{long}} = 2 l \sin \theta / \pi d$, and for $\theta = 45^\circ$, this expression reduces to

$$\frac{\varphi_{\text{circ}}}{\varphi_{\text{long}}} = \frac{l/d}{2.22}. \quad (1)$$

Thus, for example, if the storage length on a 3-mil wire is 100 mils, then a 15:1 gain in flux change (or voltage) is obtained.

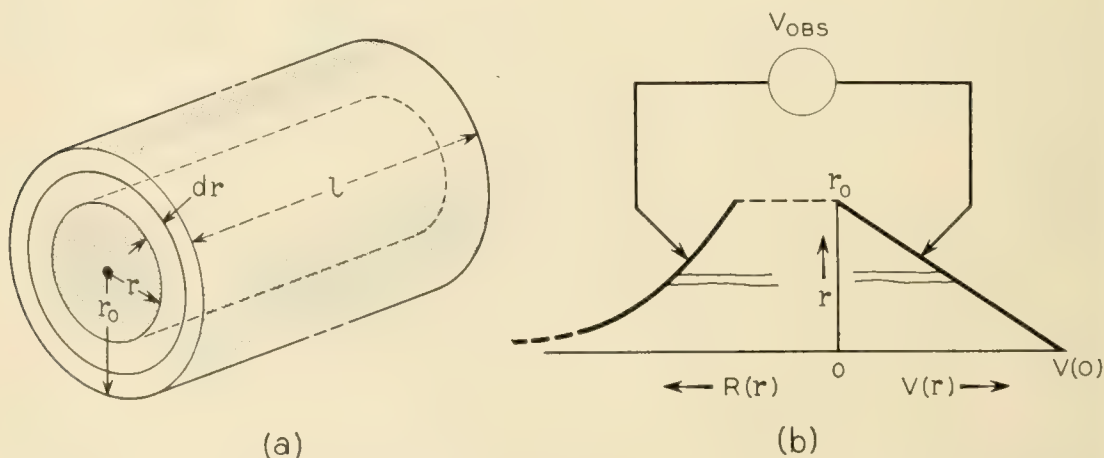


Fig. 7 — (a) Calculation of the observable voltage V_{obs} for a solid magnetic wire. (b) Diagram of induced voltage $V(r)$ and resistance $R(r)$.

The above derivation assumed that the entire circular flux change could be observed externally. Since the magnetic wire must, of necessity, serve as the source of the generated voltage the resultant eddy current flow reduces the *observable* flux change by a factor of three. Consider Fig. 7; assume that the flux reversal takes place in the classical manner and consider the circular component of this flux since it alone contributes to the observable signal. The induced voltage $V(r)$ at any point r is $V(r) = V(0)/(r_0 - r/r_0)$, where r_0 is the radius of the wire and $V(0)$ the voltage at $r = 0$. But $V(r)$, the induced or open-circuit voltage per length of wire l , could only exist if the wire were composed of many concentric tubes of wall thickness dr , each insulated from one another. In a long wire no radial eddy currents can exist. Therefore the wire of length l can be assumed to be faced by a perfect conductor at both of its ends. It remains to calculate the potential between these ends. The resistance of the tube is $R(r) = \rho l / 2\pi r dr$, where ρ is the resistivity in ohm-cm. The resistance of the wire is given by $R_0 = \rho l / \pi r_0^2$. These resistances form a voltage divider on the *induced* voltage in the tube and the total contribution of all tubes is obtained by integration;

$$\begin{aligned} V_{\text{observed}} &= \int_0^{r_0} V(r) \left(\frac{\rho l / \pi r_0^2}{\rho l / 2\pi r dr} \right) \\ &= \int_0^{r_0} V(0) \left(\frac{r_0 - r}{r_0} \right) \frac{2r}{r_0^2} dr \\ &= \frac{V(0)}{3}. \end{aligned} \tag{2}$$

Thus, (1) must be modified by (2) with the voltage step-up per memory cell becoming,

$$\frac{V_{\text{obs}}}{V_{\text{long}}} = \frac{l/d}{6.66}. \tag{3}$$

3.11 Bulk Flux Reversal — Classical Case

The switching performance of a magnetic wire under transient conditions will now be considered. The speed of magnetization reversal of magnetic materials under pulse conditions is best characterized by s_w , the switching coefficient, usually expressed in oersted-microseconds. It is defined as the reciprocal of the slope of the $1/T_s$ versus H curve where T_s is the time required to reverse the magnetization state and H is the applied magnetic field intensity. Only eddy current losses will be considered.

First, consider the case in which the magnetization is entirely circular. Reversal from one flux remanent state to the other is assumed to occur uniformly in time T_s . Axial eddy currents will flow down the center of the wire and return along the surface. The switching coefficient, s_w , will be obtained by equating the input energy to the dissipated energy. The total energy dissipated per unit length is,

$$\mathcal{E} = T_s \int_0^{r_0} 2\pi r P(r) dr, \quad (4)$$

and

$$P(r) = \frac{[E(r)]^2}{\rho}, \quad (5)$$

where $P(r)$ the power density is given by $E(r)$ the voltage gradient squared divided by the resistivity. The average energy per unit volume is therefore

$$\begin{aligned} \mathcal{E}_{av/cm^3} &= \frac{T_s}{\pi r_0^2} \int_0^{r_0} \frac{\left(V(0) \left(\frac{r_0 - r}{r_0} \right) - \frac{V(0)}{3} \right)^2}{\rho l^2} 2\pi r dr \\ &= \frac{T_s}{18\rho} \left(\frac{V(0)}{l} \right)^2. \end{aligned} \quad (6)$$

Now $V(0) = [(dB/dt)r_0 l]10^{-8}$, so $V(0)/l = (2B_s r_0/T_s)10^{-8}$. Putting this expression into (6) yields

$$\mathcal{E}_{aver/cm^3} = \frac{2B_s^2 r_0^2}{9\rho T_s} 10^{-16}. \quad (7)$$

The input energy per unit volume is

$$\mathcal{E}/cm^3 = \frac{(2B_s H \cos \theta_1) 10^{-7}}{4\pi}, \quad (8)$$

since $\Delta \vec{B} \cdot \vec{H} = 2B_s H \cos \theta_1$ where θ_1 is the angle between the applied field H and the switching flux. The factor $10^{-7}/4\pi$ is a constant relating the energy in joules to the BH product in gauss-oersteds. By equating (7) and (8), and replacing H by $(H - H_0)$, the desired s_w expression is obtained;

$$s_w = T_s (H - H_0) = \frac{(4\pi B_s r_0^2) 10^{-3}}{9\rho \cos \theta_1} (\text{oe-}\mu\text{sec}). \quad (9)$$

The substitution of $H - H_0$ for H requires some explanation. The switching curve of $1/T_s$ versus H is not a straight line as would be pre-

dicted from (9) but generally possesses considerable curvature at low drives. Equation (9) satisfactorily predicts the slope of the switching curve in the high drive region, but H_0 must be determined experimentally. In Section 3.12, flux reversal by wall motion is treated as it is a possible switching mechanism at low drives.

The switching coefficient s_w for the case in which the magnetization is purely axial will now be treated. As above, the flux density will change from $-B_s$ to $+B_s$ uniformly in time T_s . The eddy currents, which are circular, result from an induced voltage $V(r)$ where $V(r) = [V(r_0)(r/r_0)^2]$, and $V(r_0)$ is given by $V(r_0) = [(2B_s/T_s)\pi r_0^2]10^{-8}$. Thus, $E(r) = V(r)/2\pi r$, and $E(r) = (B_s r/T_s)10^{-8}$. Following the procedure used above, the internally dissipated energy density is

$$\begin{aligned}\mathcal{E}_{av/cm^3} &= \frac{T_s}{\pi r_0^2} \int_0^{r_0} \frac{(B_s r 10^{-8}/T_s)^2}{\rho} 2\pi r dr, \\ \mathcal{E}_{av/cm^3} &= \left(\frac{B_s^2 r_0^2}{2T_s \rho} \right) 10^{-16}.\end{aligned}\tag{10}$$

Equating this expression to (8) yields

$$s_w = (H - H_0)T_s = \frac{\pi B_s r_0^2 10^{-3}}{\rho \cos \theta_2},\tag{11}$$

where θ_2 is defined as the angle between the applied field and the switching flux now assumed axial.

The helical flux vector in a twistor can be resolved into a circular and an axial component. Fortunately, since the dissipated energy is proportional to the eddy current density squared, and the axial and circular current density vectors are perpendicular to each other, it is possible to write

$$\mathcal{E}_{av/cm^3}(\text{helical}) = \mathcal{E}_{av/cm^3}(\text{axial}) + \mathcal{E}_{av/cm^3}(\text{circular}).\tag{12}$$

It follows, for a 45 degree pitch angle, that

$$s_w(\text{helical}) = \frac{s_w(\text{axial}) + s_w(\text{circular})}{2}\tag{13}$$

where the factor "2" is a consequence of the flux density components being smaller by $1/\sqrt{2}$ than their resultant. Substitution of (9) and (11) into (13) gives the desired switching coefficient

$$s_w(\text{helical}) = (H - H_0)T_s = \frac{13\pi}{18} \frac{B_s r_0^2 10^{-3}}{\rho \cos \theta}.\tag{14}$$

The term $\cos \theta$ requires further explanation. The magnetization vector is constrained by energy considerations to align with the easy direction of magnetization. The angle between the applied field and the easy direction of magnetization is called θ . Equation (14) is valid for any direction of applied field. The angles θ_1 and θ_2 used in deriving (9) and (11) respectively are each 45 degrees for the helical pitch angle assumed above.

Equation (14) indicates that for maximum switching speed a material with low saturation flux density and high resistivity is required. The lower limit on s_w will be determined by internal loss mechanisms not treated here. Experimentally, this lower bound is found to be approximately 0.2 oe- μ sec.

3.12 Reversal by Single Wall

The switching time of a twistor when operating in a memory array under coincident current conditions will depend upon the low-drive switching coefficient. Experimentally, it is observed that the low drive s_w is several times the high-drive value. In this section, following the method of Williams, Shockley, and Kittel,⁵ flux reversal by the movement of a single wall will be treated. Only the circular flux case will be considered.

The technique used to obtain s_w is identical to that used in Section 3.1.1 except it is postulated that a single wall concentric to the wire moves either from the wire surface inward, or from the wire axis outward. The result is independent of the direction in which the wall moves. Assume the wall moves from $r = 0$ to $r = r_0$, as indicated in Fig. 8. In-

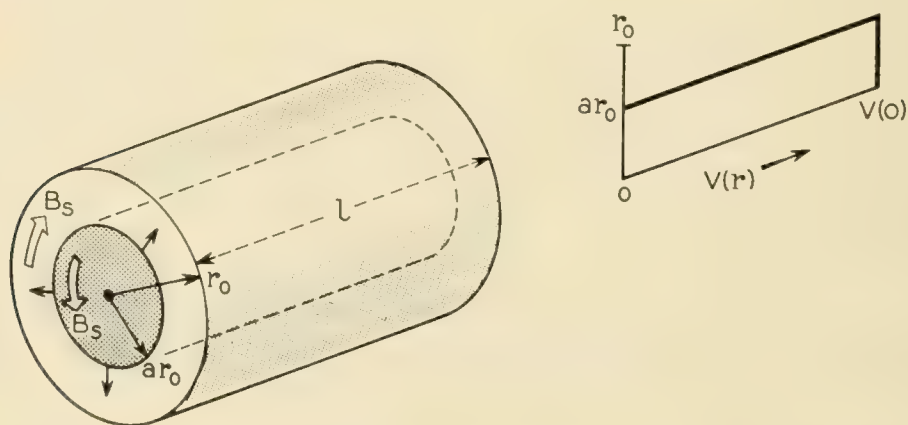


Fig. 8 — Flux reversal by expanding wall instantaneously located at radial position ar_0 moving with velocity v .

stantaneously, the wall is located at ar_0 and is traveling with a velocity v . The induced voltage $V_{oc}(r)$ is,

$$\begin{aligned} V_{oc}(r) &= (2B_s l v) 10^{-8} & 0 < r < ar_0, \\ &= 0 & ar_0 < r < r_0, \end{aligned} \quad (15)$$

and the observable voltage for a wire of length l is given by

$$\begin{aligned} V_{sc} &= \int_0^{r_0} V_{oc}(r) \frac{\rho l / \pi r_0^2}{\rho l / 2\pi r dr} \\ &= \int_0^{ar_0} V_{oc}(r) \frac{2r}{r_0^2} dr \\ &= a^2 V_{oc}(r). \end{aligned}$$

It is clear in the above integration that $V_{oc}(r)$ must be treated as a constant. Using expressions (4) and (5)

$$\begin{aligned} \frac{\partial \mathcal{E}_{av/cm^3}}{\partial t} &= \frac{1}{\pi r_0^2 \rho} \int_0^{ar_0} \left(\frac{V_{oc}^2}{l} \right) (1 - a^2)^2 2\pi r dr \\ &\quad + \frac{1}{\pi r_0^2 \rho} \int_{ar_0}^{r_0} \left(\frac{V_{oc}}{l} \right)^2 (0 - a^2)^2 2\pi r dr, \quad (16) \\ \frac{\partial \mathcal{E}_{av/cm^3}}{\partial t} &= \left(\frac{V_{oc}}{l} \right)^2 \frac{(1 - a^2)a^2}{\rho}. \end{aligned}$$

The rate of applying energy is

$$\frac{\partial \mathcal{E}}{\partial t} = \left(\frac{\partial B}{\partial t} \frac{(H - H_c) \cos \theta}{4\pi} \right) 10^{-7}. \quad (17)$$

Once again hysteresis losses are not included. Since $\partial B / \partial t = 2B_s v r_0$, and V_{oc} is given by (15), the equation of (16) and (17) yields,

$$(1 - a^2)a^2 v = \frac{\rho(H - H_c) \cos \theta}{8\pi B_s r_0}.$$

Since $v = (dr/dt) = r_0 (da/dt)$,

$$\begin{aligned} \int_0^a a^2(1 - a^2) da &= \frac{\rho(H - H_c) \cos \theta}{8\pi B_s r_0^2} \int_0^t dt, \\ \frac{a^3}{3} - \frac{a^5}{5} &= \frac{\rho(H - H_c) \cos \theta}{(8\pi B_s r_0^2) 10^{-9}} (t + K). \end{aligned}$$

When a equals 0, t equals 0, so the constant of integration is zero. When a equals 1, t equals T_s , so

$$s_w = T_s(H - H_c) = \left(\frac{16\pi}{15} \frac{B_s r_0^2}{\rho \cos \theta} \right) 10^{-3} (\text{oe-}\mu\text{sec}) \quad (18)$$

Comparison to the corresponding bulk flux reversal case indicates that the wall motion mechanism is more lossy by a factor of 2.4.

3.2 The Composite Magnetic Wire

It is apparent from the switching data of Fig. 9 that for reasonably sized solid wires ($r_0 > 1$ mil) the switching coefficient s_w is unreasonably high. The typical ferrite memory toroid, for example, when used as a memory element has an s_w of 0.6 oersted-microseconds. The only possibility for high speed coincident-current operation for solid magnetic wires is that the material have a high coercive force H_c , a conclusion not consistent with the trend toward transistor driven memory systems.

By the use of a composite wire it should be possible to reduce the eddy current losses and still preserve a reasonable wire diameter. A composite wire, by definition, will consist of a non-magnetic inner wire clad with a magnetic skin. It may be fabricated by a plating or an extruding process.

The solid wire analysis of Section 3.1 is a special case of composite

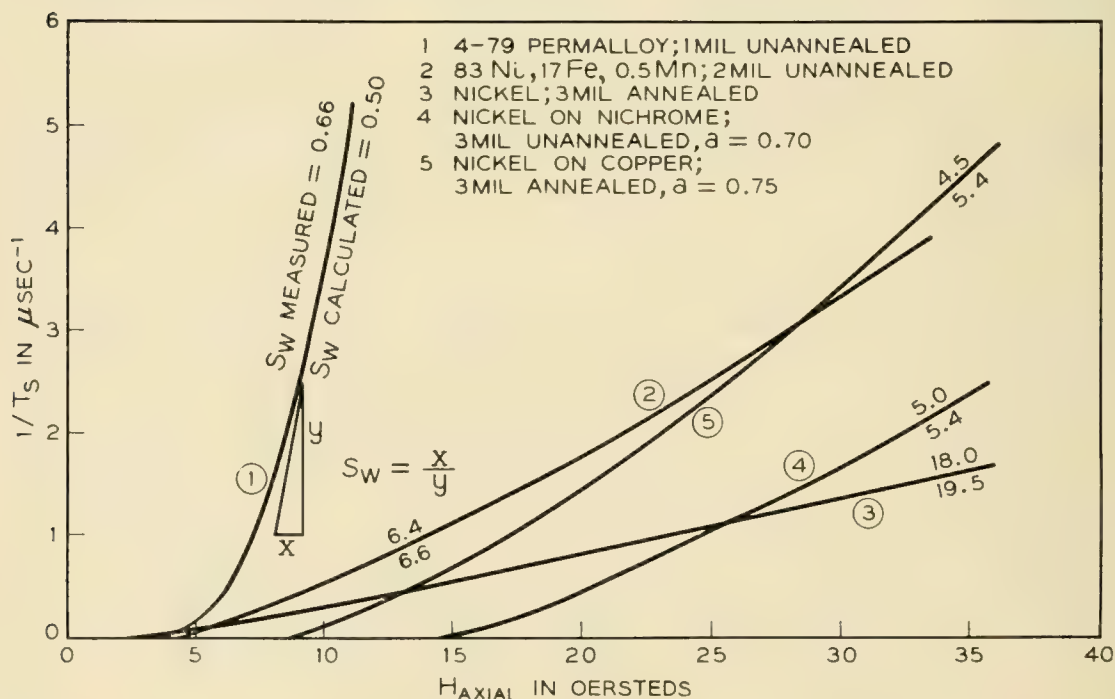


Fig. 9 — Reciprocal of flux reversal time T_s as a function of applied axial drive, H , for solid and composite magnetic wires. Sufficient torsion applied to reach saturation.

wire analysis which is given in Appendix I. Only the results of the composite wire case will be given here. As indicated in Fig. 10, ρ_1 and ρ_2

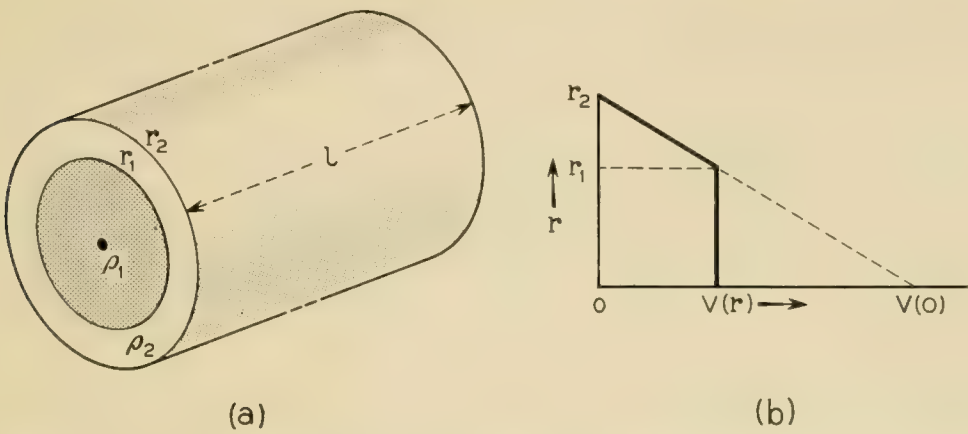


Fig. 10 — (a) Composite wire is composed of non-magnetic core covered by a magnetic skin. (b) A voltage $V(r)$ is induced in the wire during flux reversal.

are the resistivities of the inner (non-magnetic) and outer (magnetic) materials. The inner material is contained within a radius r_1 . The overall wire radius is r_2 . Defining $a = r_1/r_2$, if V_{obs} is the voltage observed across the ends of the composite wire twistor memory cell, and $V(0)$ is the induced voltage at $r = 0$ for a *solid* magnetic wire of radius r_2 , $V_{\text{obs}} = bV(0)$, and

$$b = \frac{\frac{\rho_1}{\rho_2} \left(\frac{1}{3} - a^2 + \frac{2}{3} a^3 \right) + a^2 - a^3}{\frac{\rho_1}{\rho_2} (1 - a^2) + a^2} \tag{19}$$

The parameter “ b ” reduces to $\frac{1}{3}$ for $a = 0$ in agreement with (2) which was derived for the solid wire case. Table I gives “ b ” for various material and geometry combinations.

TABLE I — THE PARAMETER “ b ”

ρ_1/ρ_2	a			
	0.9	0.8	0.7	0.0 (Solid Wire)
0	0.100	0.200	0.300	1.000
$\frac{1}{10}$	0.0988	0.194	0.285	.333
$\frac{1}{3}$	0.0963	0.184	0.259	.333
1	0.0903	0.163	0.219	.333
3	0.0790	0.135	0.180	.333
10	0.0643	0.112	0.155	.333
∞	0.0491	0.0963	0.141	.333

TABLE II — THE PARAMETER “ C_{circ} ”

ρ_1/ρ_2	a			
	0.9	0.8	0.7	0.0 (Solid Wire)
0	0.0858	0.353	0.820	1.396
$\frac{1}{10}$	0.0847	0.339	0.763	1.396
$\frac{1}{3}$	0.0842	0.311	0.657	1.396
1	0.0737	0.256	0.498	1.396
3	0.0595	0.159	0.341	1.396
10	0.0286	0.124	0.243	1.396
∞	0.0209	0.0833	0.186	1.396

The switching coefficient, s_w , for circular flux reversal, is derived as

$$s_w = \frac{(8\pi B_s r_2^2)10^{-3}}{\rho_2(1 - a^2) \cos \theta} \left\{ a^2 \left[(1 - a - b)^2 \frac{\rho_2}{\rho_1} - (1 - b)^2 + \frac{4a(1 - b)}{3} - \frac{a^2}{2} \right] + (1 - b)^2 - \frac{4(1 - b)}{3} + \frac{1}{2} \right\}, \tag{20}$$

or

$$s_w = C_{\text{circ}} \frac{(B_s r_2^2)10^{-3}}{\rho_2 \cos \theta} \text{ (oe-}\mu\text{sec)}. \tag{21}$$

Table II gives C_{circ} as a function of a and ρ_1/ρ_2 .

The switching coefficient s_w for axial flux reversal is derived as

$$s_w = \left(\frac{\pi B_s r_2^2 10^{-3}}{\rho_2 \cos \theta} \right) \left(\frac{1 - 4a^2 + a^4(3 - 4 \ln a)}{1 - a^2} \right) \tag{22}$$

$$= C_{\text{axial}} \left(\frac{B_s r_2^2 10^{-3}}{\rho_2 \cos \theta} \right). \tag{23}$$

Table III gives C_{axial} as a function of “ a ”.

Since, as explained for the solid wire case, the eddy current density vectors for circular and axial flux reversal are in quadrature,

$$s_w(\text{helical}) = \frac{s_w(\text{axial}) + s_w(\text{circular})}{2}. \tag{24}$$

Substitution of (21) and (23) into (24) gives the required expression;

$$s_w(\text{helical}) = \left(\frac{C_{\text{circ}} + C_{\text{axial}}}{2} \right) \frac{(B_s r_2^2)10^{-3}}{\rho_2 \cos \theta}. \tag{25}$$

A number of composite wire samples have been prepared and evaluated. These include nickel on nichrome and nickel on copper. The switch-

TABLE III — THE PARAMETER “ C_{axial} ”

a	0.9	0.8	0.7	0.0 (Solid Wire)
C_{axial}	0.0795	0.300	0.633	3.142

ing curves for these samples as well as for a number of solid magnetic wires are shown in Fig. 9. The agreement between the measured values of s_w and the calculated values [(14) and (25)] is quite good. Improved composite wire samples are under development.

IV. EXPERIMENTAL MEMORY CELLS AND ARRAYS

The initial experiments were performed using commercially available nickel wire of 3-mil diameter. The ϕ -NI characteristic of this wire in the helical direction is extremely square. This is a feature of *all* the magnetic materials tested whether annealed or unannealed. As a typical example,

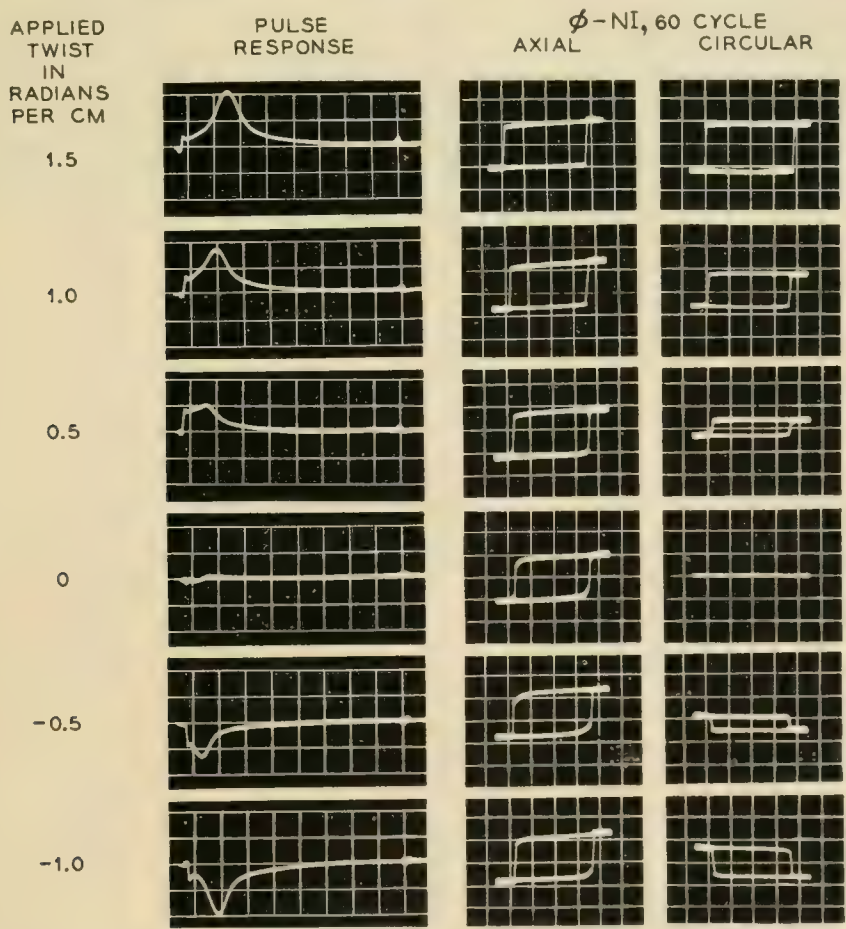


Fig. 11 — Sixty cycle and switching waveforms for 83 Ni, 17 Fe wire (see Fig. 9) as a function of applied torsion.

the 60-cycle characteristics of the axial and circular flux versus axial drive and the switching voltage waveforms under pulse conditions are given in Fig. 11 for 2-mil wire of composition 83 Ni, 17 Fe.

Note the negative prespike on the switching waveforms. By simultaneous observation of both the axial and circular switching voltage waveforms on many different magnetic wires it has been concluded that the negative prespike is due to an initial coherent rotation of the magnetization vector which results in an initial *increase* in the circular flux component. It is during this coherent rotation that the normal positive prespike on the axial switching voltage waveform is observed. Because of the mechanically introduced strain anisotropy, however, the magnetization vector is constrained to remain nearly parallel to the easy direction of magnetization. Thus, the coherent rotation soon ceases and the re-

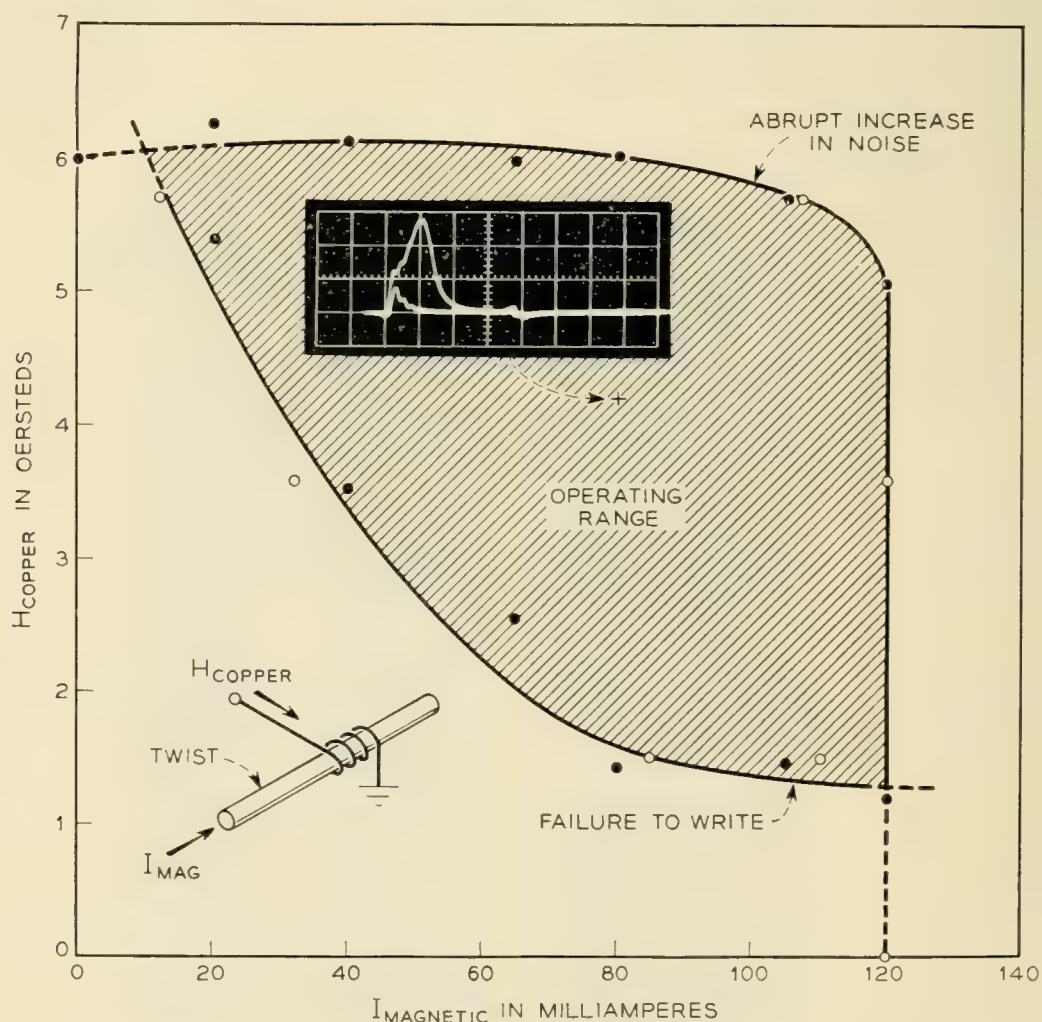


Fig. 12 — Range of writing currents for 83 Ni, 17 Fe wire operated mode. A Read drive held constant at 9 oersteds. Typical signal-to-noise ratio for a read is indicated.

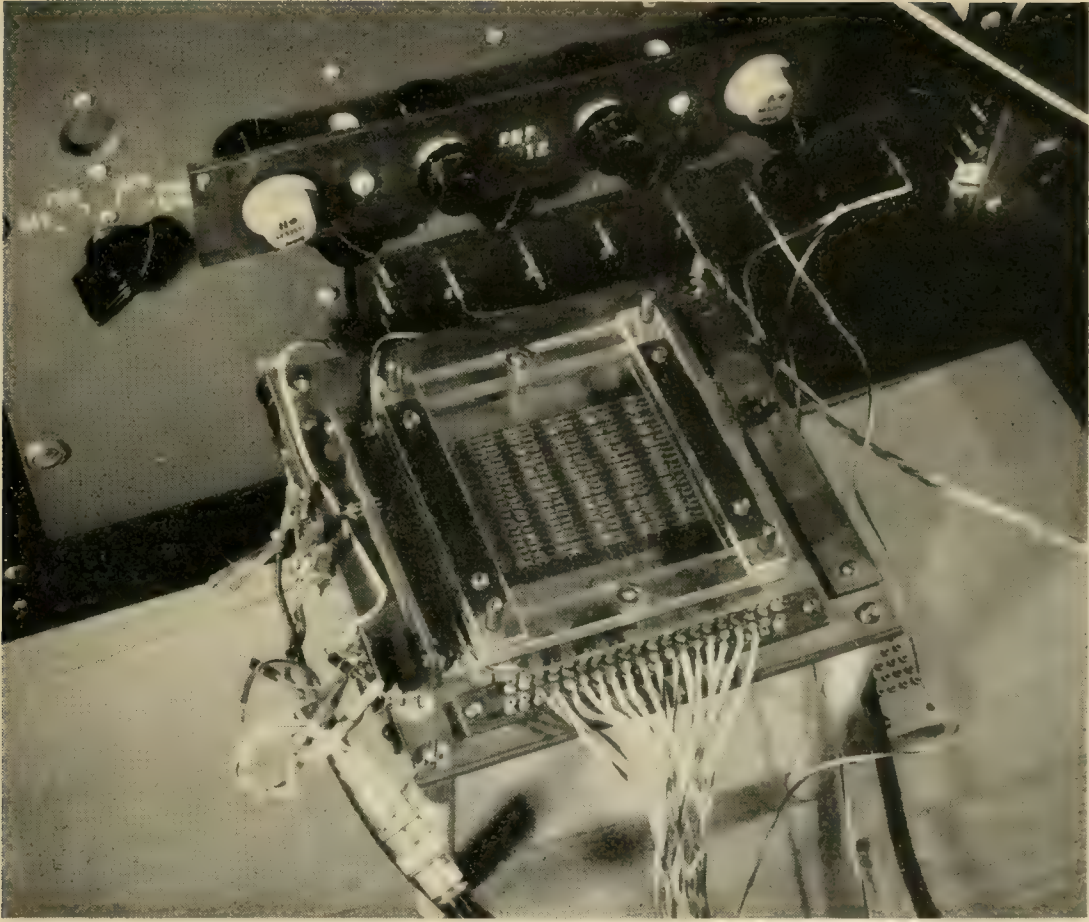


Fig. 13 — A 320-bit experimental twistor memory array. The array is transistor driven.

mainder of the flux reversal process is by an incoherent rotational process. During this latter time the circular and axial voltage waveforms are virtually identical.

Fig. 12 gives the range of operation of 2-mil 83 Ni, 17 Fe wire as a twistor operated by mode A. As a result of the extreme squareness of the ϕ -NI characteristic in helical direction the range of operation encloses an area nearly the theoretical maximum. The switching times of other memory cells tested ranged from 0.2 μ sec for a 1 mil 4-79 moly-permalloy wire to 20 μ sec for a 5 mil perminvar wire. Thus it is seen that the switching speeds of the twistor compare quite favorably with those of conventional ferrite toroids and sheets.

It is, of course, possible to store many bits of information along a single magnetic wire. The allowable number of bits per inch is related to the coercive force, the saturation flux density, and the diameter of the wire. For the nickel wire, about 10 bits per inch are possible. Predictions as to the storage density for a given material can be made by referring to

suitable demagnetization data. There are, however, interference effects between cells which are not completely understood at the present time.

A memory array (16×20) has been constructed as a test vehicle. An illustration of this array is shown in Fig. 13. The drive wires have been woven over glass tubes which house the removable magnetic wires. Provision is made for varying the torsion and the tension of the individual magnetic wires.

As an indication of the performance of the twistor, Fig. 14 is a composite photograph showing the minimum and maximum signal over the 16 bits of a given column for 3-mil nickel wire. Also included are the noise pulses for these cells, the so-called disturbed zero signals. The write currents were 2.3 ampere-turns on the solenoid and 130 ma through the magnetic wire. The read current was 6.0 ampere-turns. The array

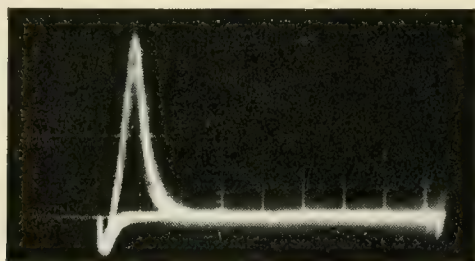


Fig. 14 — Composite photograph of the 16 output signals from a column of the array of Fig. 13. Average output signal about 3.5 millivolts; sweep speed equals $2 \mu\text{sec/cm}$.

was transistor driven. A read-write cycle time of 10 microseconds appeared to be possible.

V. DISCUSSION

The twistor is presented as a logical companion to the coincident-current ferrite core and sheet.^{6, 7} In many applications it should compete directly with its ferrite equivalents. Perhaps its greatest use will be found in very large ($>10^6$) memory arrays.

From a cost per bit viewpoint the future of the twistor appears quite promising. Fabricating and testing the wire should present no special problems as it is especially suited for rapid, automatic handling. The possibility of applying weaving techniques to the construction of a twistor matrix looks promising.

It is possible that, for both mode A and C operation of the twistor, an array can be built which consists simply of horizontal copper wires and

vertical magnetic wires — much like a window screen. Preliminary experiments have shown that single cross wires do operate successfully. The operation of this array would be analogous to a core memory array. Physically it could look just like a core array — but without the cores.

ACKNOWLEDGEMENTS

The author wishes to acknowledge the help of R. S. Title in obtaining the memory array data, R. A. Jensen in constructing the test jigs, and D. H. Wenny, Jr. in supplying the wire samples. The discussions with my associates, in particular D. H. Looney, R. S. Title, and J. A. Baldwin, have been very helpful. The writer would like to acknowledge the interest and encouragement shown by R. C. Fletcher.

APPENDIX I

From Fig. 10, for bulk circular flux reversal in a composite wire, the induced voltage $V(r)$ for a wire length l is

$$\begin{aligned} V(r) &= \left[(r_2 - r_1)l \left(\frac{2B_s}{T_s} \right) \right] 10^{-8}, & 0 < r < r_1, \\ &= \left[(r_2 - r)l \left(\frac{2B_s}{T_s} \right) \right] 10^{-8}, & r_1 < r < r_2. \end{aligned} \quad (26)$$

For a *solid* magnetic wire of radius r_2

$$V(0) = \left(\frac{2B_s r_2 l}{T_s} \right) 10^{-8}. \quad (27)$$

Therefore,

$$\begin{aligned} V(r) &= \left(\frac{r_2 - r_1}{r_2} \right) V(0), & 0 < r < r_1, \\ V(r) &= \left(\frac{r_2 - r}{r_2} \right) V(0), & r_1 < r < r_2. \end{aligned} \quad (28)$$

In general, the resistance of a tube of wall thickness dr is $R(r) = \rho l /$

$2\pi r dr$. The resistance of the wire is $R_T = \rho_1 \rho_2 l / \pi [\rho_1 (r_2^2 - r_1^2) + \rho_2 r_1^2]$. The observable voltage for a length of wire, l , is

$$\begin{aligned} V_{\text{obs}} &= \int [V(r)] \quad (\text{Volt-Divider}) \\ &= \int_0^{r_1} \left(\frac{r_2 - r_1}{r_2} \right) V(0) \left(\frac{\frac{\rho_1 \rho_2 l}{\pi [\rho_1 (r_2^2 - r_1^2) + \rho_2 r_1^2]}}{\frac{\rho_1 l}{2\pi r dr}} \right) \\ &\quad + \int_{r_1}^{r_2} \left(\frac{r_2 - r}{r_2} \right) V(0) \left(\frac{\frac{\rho_1 \rho_2 l}{\pi [\rho_1 (r_2^2 - r_1^2) + \rho_2 r_1^2]}}{\frac{\rho_2 l}{2\pi r dr}} \right). \end{aligned}$$

This reduces to

$$V_{\text{obs}} = V(0) \frac{(\frac{1}{3} - a^2 + \frac{2}{3}a^3)\rho_1/\rho_2 + a^2 - a^3}{(1 - a^2)\rho_1/\rho_2 + a^2} \quad (29)$$

$$= bV(0), \quad (30)$$

where $a = r_1/r_2$ and b is given by reference to (29). The ratio $b/\frac{1}{3} = 3b$ is the relative efficiency of the composite as compared to the solid magnetic wire from an available signal viewpoint. An expression for s_w will now be derived.

The total energy dissipated per unit length l is

$$\xi_T/l = T_s \int_0^{r_2} \rho i_d^2(r) 2\pi r dr, \quad (31)$$

where $i_d(r)$ is the current density. Now, $i_d(r) = V(r) - V_{\text{obs}}/\rho l$, therefore

$$\begin{aligned} i_d(r) &= (1 - a - b) \frac{V(0)}{\rho_1 l}, & 0 < r < r_1, \\ &= \left(1 - \frac{r}{r_2} - b \right) \frac{V(0)}{\rho_2 l} & r_1 < r < r_2. \end{aligned} \quad (32)$$

The substitution of (32) into (31) yields, after manipulations,

$$\begin{aligned} \xi_{\text{av}}/l &= \frac{\pi T_s r_2^2}{\rho_2} \left(\frac{V(0)}{l} \right)^2 \left\{ a^2 \left[(1 - a - b)^2 \frac{\rho_2}{\rho_1} - (1 - b)^2 \right. \right. \\ &\quad \left. \left. + \frac{4a(1 - b)}{3} - \frac{a^2}{2} \right] + (1 - b)^2 - \frac{4(1 - b)}{3} + \frac{1}{2} \right\}. \end{aligned} \quad (33)$$

From (8), the applied energy per wire length l is

$$\xi/l = \left(\frac{[B_s(H - H_0) \cos \theta] 10^{-7}}{2\pi} \right) \pi(r_2^2 - r_1^2), \quad (34)$$

where only that part of the applied energy associated with the high drive dynamic losses is included. Equating (33) and (34) and replacing $V(0)/l$ by (27) results in

$$s_w = (H - H_0)T_s = \frac{8\pi B_s r_2^2 10^{-3}}{\rho_2(1 - a^2) \cos \theta} \left\{ a^2 \left[(1 - a - b)^2 \frac{\rho_2}{\rho_1} - (1 - b)^2 + \frac{4a(1 - b)}{3} - \frac{a^2}{2} \right] + (1 - b)^2 - \frac{4(1 - b)}{3} + \frac{1}{2} \right\}. \quad (35)$$

This can be expressed as

$$s_w = (H - H_0)T_s = C_{\text{circ}} \frac{(B_s r_2^2) 10^{-3}}{\rho_2 \cos \theta} \text{ (oe-}\mu\text{sec)}. \quad (36)$$

Bulk axial flux reversal in a composite magnetic wire can be treated in a manner analogous to that used in Section 3.1.1 for the solid wire. The uniform reversal of the axial flux induces a voltage $V(r)$ in the wire where

$$\begin{aligned} V(r) &= V(r_2) \left[\left(\frac{r}{r_2} \right)^2 - \left(\frac{r_1}{r_2} \right)^2 \right], & r_1 < r < r_2, \\ &= 0, & 0 < r < r_1, \end{aligned}$$

and $V(r_2) = [(2B_s/T_s)\pi r_2^2] 10^{-8}$. Since $E(r) = V(r)/2\pi r$,

$$E(r) = \frac{B_s}{T_s} \left(r - \frac{r_1^2}{r} \right) 10^{-8}. \quad (37)$$

Following the procedure of Section 3.1.1.,

$$\begin{aligned} \xi_{\text{av/cm}^3} &= \frac{T_s 10^{-16}}{\pi(r_2^2 - r_1^2)} \int_{r_1}^{r_2} \frac{B_s^2}{\rho_2 T_s^2} \left(r - \frac{r_1^2}{r} \right)^2 2\pi r \, dr \\ &= \left\{ \frac{2B_s^2 r_2^2 10^{-16}}{\rho_2 T_s} \left[\frac{\frac{1}{4} - a^2 + a^4(\frac{3}{4} - \ln a)}{1 - a^2} \right] \right\}, \end{aligned} \quad (38)$$

where $a = r_1/r_2$ as before. Equating this expression to (8) yields

$$s_w = (H - H_0)T_s = \frac{\pi B_s r_2^2 10^{-3}}{\rho_2 \cos \theta} \left[\frac{1 - 4a^2 + a^4(3 - 4 \ln a)}{1 - a^2} \right] \quad (39)$$

$$= C_{\text{axial}} \left(\frac{B_s r_2^2 10^{-3}}{\rho_2 \cos \theta} \right) \text{ (oe-}\mu\text{sec)}. \quad (40)$$

REFERENCES

1. R. M. Bozorth, *Ferromagnetism*, D. Van Nostrand Company, Inc., New York, N. Y., 1951, p. 628.
2. U. F. Gianola, Use of Wiedemann Effect for Magnetostrictive Coupling of Crossed Coils, *T. Appl. Phys.*, **26**, Sept. 1955, pp. 1152-1157.
3. H. F. Girvin, *Strength of Materials*, International Textbook Co., Scranton, Pa., 1944, p. 233.
4. J. A. Baldwin, unpublished report.
5. H. J. Williams, W. Shockley, and C. Kittel, Velocity of a Ferromagnetic Domain Boundary, *Phys. Rev.*, **80**, Dec. 1950, pp. 1090-1094.
6. J. A. Rajchman, Ferrite Apertured Plate for Random Access Memory, *Proc. I.R.E.*, **45**, (March, 1957), pp. 325-334.
7. R. H. Meinken, A Memory Array in a Sheet of Ferrite, presented at Conference on Magnetism and Magnetic Materials, Boston, Mass., Oct. 16-18, 1956.

Non-Binary Error Correction Codes*

By WERNER ULRICH

(Manuscript received April 19, 1957)

If a noisy channel is used to transmit more than two distinct signals, information may have to be specially coded to permit occasional errors to be corrected. If pulse amplitude modulation is used, the most probable error is a small one, e.g., 6 is changed to 7 or 5. Codes for correcting single small errors, and for correcting single small errors and detecting double small errors, in a message of arbitrary length, for an arbitrary number of different signals in the channel, are derived in this paper.

For more specialized situations, the error is not necessarily restricted to a small value. Codes are derived for correcting any single unrestricted error in a message of arbitrary length for an arbitrary number of different signals.

Finally, a set of codes based partially upon the Reed-Muller codes is described for correcting a number of errors in a more restricted class of message lengths for an arbitrary number of different signals.

The described codes are readily implemented. Many techniques are used which have an analog in a binary system. Other techniques are broadly analogous to binary coding techniques or are special adaptations of a binary code.

1. INTRODUCTION

1.1 Use of Error Correction Codes

One function of an error correction code is to aid in the correct transmission of digital information over a noisy channel. This process is illustrated in Fig. 1. An information source gives information to an encoder; the encoder converts the information into a message containing sufficient redundancy to permit the message to be slightly mutilated by the noisy channel and still be correctly interpreted at the destination. The message is then sent via the noisy channel to a decoder which will

* This paper was submitted to Columbia University in partial fulfillment of the requirements for the degree of Doctor of Engineering Science in the Faculty of Engineering.

reconstruct the original information if the mutilation has not been excessive. Finally, the information is sent to an information receptor.

One scheme for correcting errors in a binary system is to send each binary digit of information three times and to accept at the receiver that value which is represented by two or three of the received digits. Then, the encoder is simply an instrument for causing each digit to be sent three times, and the decoder consists of a majority organ. However, many methods are available which are considerably more elegant, and which will permit more information to be passed through a noisy channel in a given unit of time. This paper will deal with such methods for channels capable of sending b different symbols instead of the usual 1 and 0 of a binary channel.

The most convenient explanation of an error correction code has been made with respect to the transmission of correct digital information over a noisy channel. This does not imply the restriction of such codes

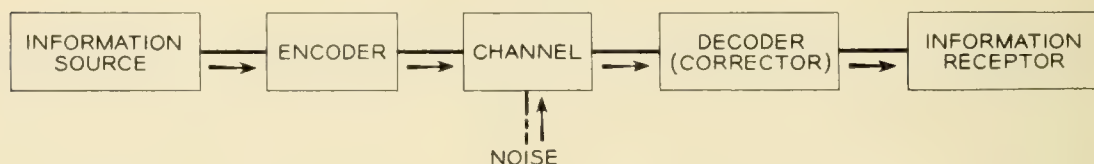


Fig. 1 — Transmission over a noisy channel.

to the noisy channel problem exclusively. Actually, the first application considered for such a code was with respect to computers.¹ Many large high speed computers stop whenever an error is detected in some calculation and must be restarted; with the use of an error correction code this could be avoided by permitting the computer to correct its own random errors directly. To the best knowledge of the author, error correction codes have not yet been used in any major computer. But the storage system of a computer may, in the future, lend itself to the use of error correction codes.

Frequently, very elaborate precautions must be taken in present storage systems to insure that they are free from errors. Magnetic tapes must be specially made and handled to guarantee the absence of defects, magnetic cores must be carefully tested to make sure that no defective cores get into an array, cathode ray tubes used in Williams Tube or Barrier Grid Tube storage systems must be perfect. Probably, there are other storage methods whose development is hampered because of a common requirement for error-free performance in all storage locations. With the use of error correction codes, such storage systems could be used, if they are sufficiently close to perfection, even though not perfect.

It is not unlikely that the near future will see the development of storage systems which will be able to store more than two states at every basic storage location.² If such systems are developed, it seems likely that they will be more erratic or noisy than binary storage systems, since each location must store one of b signals instead of one of two. If a cathode ray tube storage system were used, for example, different quantities of charge would have to be distinguished; in a binary storage system, only the presence or absence of charge must be detected. This suggests that error correction codes may become essential with certain types of non-binary storage systems. One object of this paper is to develop codes for this purpose and to discover which number systems are most easily correctable.

Some investigations have been made on the use of computer systems using multi-state elements.³ A switching algebra has been developed similar to Boolean algebra for handling switching problems in terms of multi-state elements. Single device ring counters (the cold cathode gas stepping tube for example) already exist and might be useful in such systems. But currently, only limited steps in this direction have been made. Another object of this paper is to show the advantages and problems of error correction codes in multi-state systems; it is not unreasonable to predict that error correction codes may be more necessary in multi-state systems than in binary systems.

1.2 *Geometric Concept of Error Correction Codes*

A geometric model of a code was suggested by R. W. Hamming¹ which can be altered slightly to fit the non-binary case. For an n digit message, a particular message is a point in n dimensional space. A single error, however defined, will change the message, and will correspond to another point in n dimensional space. The distance between the original point and the new point is considered to be unity. Thus, the distance d between the points corresponding to any two messages is defined as the minimum number of errors which can convert the first message into the second.

With an error detection and/or correction code, the set of transmitted messages is limited so that those which are correctly received are recognizable; those messages which are received with fewer than a given number of errors are either corrected or the fact that they are wrong is recognized and some other appropriate action (such as stopping a computer) is taken.

In the case of binary codes, an error changes a 1 to 0 or a 0 to 1. In the non-binary case, two definitions of an error are possible and will be

used in this paper. A *small* error changes a digit to an adjacent value. In a decimal system, a change from 1 to 2 or 1 to 0 is a small error. An *unrestricted* error changes a digit to any other value. In a decimal system, a change from 1 to 5 is an unrestricted error.

1.3 *Material To Be Presented*

The various types of codes described in this paper and the sections in which they are to be found are summarized in Table I. The techniques which are described are summarized below.

The geometric model suggests the simplest approach to error correction codes. A transmitter has a “codebook” containing all members of the set of transmitted messages. If the message source gives to the encoder the signal that the information to be sent is k (that is to say, the k th

TABLE I — TYPES OF CODES

Type of Code	Distance	Type of Error	Described in Section
Single Error Detection	2	Small and Unrestricted	II
Single Error Correction	3	Small	III and 6.1
Single Error Correction Prime Number Base	3	Unrestricted	4.1
Composite Number Base		Unrestricted	4.2
Single Error Correction and	4	Small	V and 6.1
Double Error Detection			
Multiple Error Correction	—	Small	6.2

output of all the outputs associated with the message source), the encoder chooses the k th member of the set. The decoder will then look up the message it receives in its own codebook which contains all possible received messages, and corresponding to the entry of the received message will find the symbols corresponding to k . Or the receiver may compare the received message with every member of the set of transmitted messages, calculate the distance between the two, and correct the received message to whichever of the transmitted messages is separated from the received message by the smallest distance. (It has been shown by Slepian⁴ that this is the message most likely to be correct in a symmetrical binary channel having the property that changes from 1 to 0 and from 0 to 1 as a result of noise in the channel are equally likely.)

The practical difficulty with such a code is the large size of the required codebooks. Most coding schemes try to eliminate such codebooks and substitute a set of rules for encoding, decoding and correcting messages.

One approach toward creating a simple association between the information and the message is to use some of the digits of the message for conveying information directly. The Hamming Code¹ uses this technique.

An *information digit* is a digit of a message that is produced directly by the information source; in a base b code, an information digit may have b different values, the choice between these values representing the information that is to be sent.

A *check digit* is a digit of a message that is calculated as a function of the information digits by the encoder. It is sometimes convenient to represent or calculate a check digit in terms of a recursive formula using previously calculated check digits as well as information digits. In a base b code, a check digit may have b *check states*. When more than one check digit is used, each different combination of check digits corresponds to a different check state for the message; a message with m check digits will have b^m *message check states*.

A *systematic*⁵ code encoder generates messages containing only information digits and check digits. The information source generates only base b information digits. The Hamming Code is a systematic code.

Section II offers a general method for obtaining single error detection codes for both small and unrestricted errors. The idea of mixed digits (digits which are, in a sense, neither information nor check digits, but a combination of both) is introduced, and it is shown how mixed digits may lead to more efficient coding systems. This idea is believed to be novel. Code systems which use mixed digits are called semi-systematic codes. Semi-systematic codes are used extensively throughout this paper.

Section III offers a general method for obtaining single small error correction codes, including both systematic and semi-systematic codes.

Section IV offers a general method for obtaining the more complicated single unrestricted error correction codes. The problem is divided into two parts. Section 4.1 describes codes for correcting single unrestricted errors in case b , the base of the channel, is a prime number.* Section 4.2 describes a special technique for obtaining the more complex codes for correcting single unrestricted errors in the event b is a composite number.

Section V offers a general method for obtaining semi-systematic codes for correcting single small errors and detecting double small errors. No general solution has been found for obtaining single error correction or double error detection codes for the case of unrestricted errors. No gen-

* This class of codes was previously described in a brief summary by Golay.⁶

eral solution has been found for multiple error correction codes for the unrestricted error case.

In Section VI, a number of techniques are presented for using binary error correction coding schemes for non-binary error correction codes. Section 6.1 shows how such techniques may be used to obtain non-binary single error correction codes, and single error correction double error detection codes, for the small error case. Section 6.2 presents a special technique, involving the use of an adaptation of the Reed-Muller binary code, to obtain a class of non-binary multiple error correction codes, for the small error case.

Section VII shows that an iterative technique of binary coding can be directly applied to non-binary codes. It also shows how an adapted Reed-Muller code can be profitably used in such a system.

Section VIII summarizes the results obtained in Sections II-VII and shows the advantages and shortcomings of many of these codes.

Section IX presents general conclusions which may be drawn from this paper.

II. SINGLE ERROR DETECTION CODES

Single error detection codes require message points separated in n dimensional space by a distance of two.

For the binary case, the only two possible types of errors are the change from a 1 to a 0 and from a 0 to a 1.

A simple technique that is used frequently for binary error detection codes is to encode all messages in such a manner that every message contains an even number of 1's. This is accomplished by adding a *parity check digit* to the information digits of a message; this digit is a 1 if an odd number of 1's exist in the information digits of a message and is a 0 if an even number of 1's exist in the information digits. At least two errors must occur before a message containing an even number of 1's can be converted into another message containing an even number of 1's, since the first error will always cause an odd number of 1's to appear. A message with an odd number of 1's is known to be incorrect.*

An analogous technique may be used for the unrestricted error case in non-binary codes. We can obtain a satisfactory code by adding a complementing digit to a series of information digits to form a message.

A *complementing digit*, base b , is defined as a digit which when added to some other digit will yield a multiple of b .

* Parity check digits may be selected to make the number of 1's in a message always odd, but the principle is the same; in this case, an error is recognized if a received message contains an even number of 1's.

For a single unrestricted error detection code, the complementing digit complements the sum of the information digits. A complementing digit is a check digit. In the binary case, it is a parity check digit.

As an example, consider a decimal code of this type. A message 823 would require a complementing digit 7, making the total message 8237 ($8 + 2 + 3 + 7 = 20$, a multiple of 10). An error in any one digit will mean that the sum of the message digits will not be a multiple of 10.

For the small error case, it is sufficient to make certain that the sum of all digits is even since any error of ± 1 would destroy this property. For the binary case, all errors are small since the only possible error on any digit is a change by ± 1 ; a simple parity check is adequate. For a non-binary code, it would be wasteful to add a digit just to make sure that the sum of all digits is even. In a decimal code for example, if the sum of the message digits is even, the values 0, 2, 4, 6, 8 for the check digit will satisfy a check, or if the sum of the message digits is odd, the values 1, 3, 5, 7, 9 will satisfy the check. More information could be sent if a choice among these values could be associated with information generated by the information source.

This introduces the concept of a mixed digit; i.e., a digit which conveys both check information and message information.

A *mixed digit* is defined as follows: a mixed digit x , base b , is composed of two components (y, z) where y represents an information component and z represents a check component. The number of information states of a mixed digit is β , with y taking the values $0, 1, \dots, \beta - 1$; the number of check states of a mixed digit is α , the number base of z . In a message containing m check digits and h mixed digits, the number of check states for the *message* is $b^m \cdot \alpha_1 \cdot \alpha_2 \cdot \dots \cdot \alpha_h$, where α_i is the number of check states of the i 'th mixed digit.

If mixed digits are used as part of a code, information must be available in at least two number bases; b , the number base of the channel, and β , the number base of the mixed digit. A situation where this arises naturally is in the case of the algebraic sign of a number; this is a digit of information, base 2, which may be associated with other digits of any base. Similarly, any identification which must be associated with numerical information can be conveniently coded in a number base different from the number base of the numerical information. Thus, a mixed digit can sometimes be used conveniently in an information transmission system without complicating the information source and receptor.

An error detection code for single small errors suggests the use of a mixed digit. In the decimal code for example, the quibinary⁷ representa-

TABLE II — QUIBINARY CODE

Quinary Component	Binary Component	Decimal Digit
0	0	0
0	1	1
1	0	2
1	1	3
2	0	4
2	1	5
3	0	6
3	1	7
4	0	8
4	1	9

tion of the mixed digit might be used, letting the quinary component of the mixed digit convey information and the binary component a check. (Table II.)

The information source generates blocks of decimal digits followed by one quinary digit. The messages are then generated in the following way: record all decimal information digits as information message digits and take their sum; if the sum is even, the binary component, z , of the mixed digit is 0, otherwise it is 1. The quinary component, y , of the mixed digit is taken directly from the information source and combined with the calculated binary part by the rules of the quibinary code to form the mixed decimal digit. Thus, x , the value of the mixed digit, is given by the formula:

$$x = 2y + z. \tag{1}$$

For example, if the decimal digits of a message are 289 and the quinary digit of the message is 3, the mixed digit is 7, and the message is 2897. The sum of the decimal information digits is 19, which is odd, so that the binary component of the mixed digit is 1; this is combined with the quinary component, 3, by the rules of the quibinary code table, to form decimal digit 7. The requirement that the sum of all digits be even is satisfied by the binary component of the mixed digit, and the information associated with the mixed digit is contained in the quinary component.

This method is easily extensible to any other number base and is also extensible to the case of slightly larger but still restricted errors (such as ± 1 or ± 2), provided that the maximum single error is less than $(b - 1)/2$.

From the preceding example, it is apparent that mixed digits can be usefully employed in error detection codes. The use of mixed, check and

information digits simplified the encoder and decoder. To differentiate among the classes of codes which will be described in this paper, the following terms will be used, in addition to those previously defined.

A *semi-systematic code* encoder produces messages containing only information, mixed and check digits. The information source generates information digits in base b for information digits, and in base β for mixed digits. (The example given above is a semi-systematic code.)

Of two coding schemes in the same channel base b , each working with messages of the same length, and each satisfying a given error detection or correction criterion, the more *efficient* scheme is defined as the one which produces the larger number of different possible messages.

III. SINGLE ERROR CORRECTION CODES, SMALL ERRORS (± 1)

The problems of error correction codes in nonbinary systems are extensive and must be treated in several distinct sections. The basic difference between the error correction problem in binary and non-binary codes is the fact that the sign of the error is important. In a binary code, if the message 11 is received and it is known that the second digit is incorrect, only one correction can be made, to 10. But in a decimal code with errors limited to ± 1 , if the message 12 is received and it is known that the second digit is wrong, it can be changed to either 11 or 13.

Consider the following simple code for correcting single small errors. A decimal channel is used, and a message is composed of three information digits and one check digit. Let x_1 represent the check digit and x_2, x_3, x_4 the information digits. Here, x_1 is chosen to satisfy*

$$x_1 + 2x_2 + 3x_3 + 4x_4 = 0 \text{ mod } 10. \quad (2)$$

The encoder calculates x_1 , and transmits the message $x_1x_2x_3x_4$. This is received as $x_1'x_2'x_3'x_4'$. The decoder then calculates c given by

$$c = (x_1' + 2x_2' + 3x_3' + 4x_4') \text{ mod } 10. \quad (3)$$

If the assumption is made that at most a single small error exists, then this error can be corrected by using the following rules, which may be verified by inspection.

If $c = 0$, no correction is necessary;

$5 > c > 0$, decrease the c th digit by one;

* By definition $a = c \text{ mod } b$ is equivalent to $a = c + nb$, where a, b, c and n are integers. The equality notation is used in preference to the congruence notation throughout this paper, since an addition performed without carry occurs naturally in many circuits; in terms of such a circuit, the $\text{mod } b$ signifies only the base of the addition, and a true equality exists between the state of two circuits, with the same output even though one has been cycled more often.

$5 < c$, increase the $(10 - c)$ th digit by one;

$c = 5$ implies a multiple error or a larger error.

Since the value of c is used for correcting a received message, it is called the corrector.* For the general case, a corrector is defined as follows.

In a message encoded to satisfy m separate checks, the result of calculating the checks for the received message at the decoder is an m digit word called the *corrector*. There are as many possible values of the corrector as there are check states of the message, although all of the values of the corrector need not correspond to a correctable error.

It is important that, for a given transmitted message, every different error will lead to a different value of the corrector; otherwise there will be no way of knowing which correction corresponds to a particular value of the corrector. The number of correctable errors may be far less than the number of possible values of the corrector, so that not all of these values may be useful for a code to correct a particular class of errors. However, the number of corrector states sets an upper limit to the number of possible corrections.

For many codes, it is convenient to associate a particular value of a corrector for the condition that a particular digit has been received too high by a single increment, for example, a 7 received as an 8.

The *characteristic* of a digit for a particular code is defined as the value of the corrector if that digit is incorrectly received, the error having increased the value of the digit by $+1$, and all other digits are correctly received. Obviously, this definition only applies to those codes having the property that the value of the corrector is independent of the value of the incorrect digit and of the other digits.

A *simple characteristic code* encoder produces messages in which each digit has a distinct characteristic as defined above.

The Hamming code is an example of a simple characteristic code as is the code previously described. In that example, the characteristic of x_i is i .

The advantage of a simple characteristic code for single small error correction is obvious: the association between the calculated checks and the correction to be performed is simple and does not depend on the values of the digits of the message.

The following example of a simple characteristic code will illustrate this principle more fully.

Consider a single small error correction code, working with a quinary

* The terms corrector and characteristic were first used in a more restricted sense in an article on binary coding by Golay.⁸

(base 5) channel. Each message will consist of ten information digits and two check digits.

Let x_1 and x_2 represent the check digits, and $x_3, x_4, \dots, x_{12}$ represent the information digits.

The equations for calculating x_1 and x_2 are:

$$\begin{aligned} 1x_1 + 0x_2 + 0x_3 + 1x_4 + 1x_5 + 1x_6 + 1x_7 \\ + 2x_8 + 2x_9 + 2x_{10} + 2x_{11} + 2x_{12} = 0 \bmod 5, \end{aligned} \quad (4)$$

$$\begin{aligned} 0x_1 + 1x_2 + 2x_3 + 1x_4 + 2x_5 + 3x_6 + 4x_7 \\ + 0x_8 + 1x_9 + 2x_{10} + 3x_{11} + 4x_{12} = 0 \bmod 5. \end{aligned} \quad (5)$$

At the decoder, the corrector terms, c_1 and c_2 , are calculated using x'_i , the received value of x_i , in the following formulas:

$$\begin{aligned} 1x'_1 + 0x'_2 + 0x'_3 + 1x'_4 + 1x'_5 + 1x'_6 + 1x'_7 \\ + 2x'_8 + 2x'_9 + 2x'_{10} + 2x'_{11} + 2x'_{12} = c_1 \bmod 5, \end{aligned} \quad (6)$$

$$\begin{aligned} 0x'_1 + 1x'_2 + 2x'_3 + 1x'_4 + 2x'_5 + 3x'_6 + 4x'_7 \\ + 0x'_8 + 1x'_9 + 2x'_{10} + 3x'_{11} + 4x'_{12} = c_2 \bmod 5. \end{aligned} \quad (7)$$

The values of c_1c_2 corresponding to the condition that one and only one digit is too high by 1, $x'_i = x_i + 1$, can be read by reading the coefficients of the i th digit in the corrector formulas. This quantity is therefore the characteristic of the i th digit. If $x'_i = x_i - 1$, then the fives complements of these coefficients will be the value of the corrector. Table III lists the characteristics and characteristic complements associated with each digit.

TABLE III — CHARACTERISTICS AND CHARACTERISTIC COMPLEMENTS
SYSTEMATIC QUINARY CODE

Digit	Characteristic	Complement of Characteristic
x_1	10	40
x_2	01	04
x_3	02	03
x_4	11	44
x_5	12	43
x_6	13	42
x_7	14	41
x_8	20	30
x_9	21	34
x_{10}	22	33
x_{11}	23	32
x_{12}	24	31

In this code all the possible values of c_1c_2 correspond to the characteristic of a digit or the complement of this characteristic, except 00 which corresponds to the correct message. (An inspection of equations (4) through (7) reveals that if $x_i' = x_i$ for all values of i , the values of c_1 and c_2 are 0). Thus, we can assign a unique correction to each value of c_1c_2 .

The above techniques are extensible to other number bases and different length words provided b , the number base of the channel, is greater than 2. (The equivalent binary channel problem has been treated by Hamming.¹) The following set of rules and conventions may be used for deriving a satisfactory set of characteristics for a simple characteristic systematic code used to correct single small errors for any length message, and any base, $b \geq 3$. The rules must be followed, and the conventions (which represent one pair of conventions out of the set of pairs of conventions, which together with Rules 1 and 2 can be used for deriving a code of this class) if followed, will lead to a reasonably simple method for encoding and decoding messages.* Since the rules, not the conventions, limit the efficiency of the code, no set of conventions can be found which will lead to a more efficient code of this class.

Rule 1. For an n digit message (including check digits), m check digits are required and m must satisfy the following inequalities:

$$\text{if } b \text{ is odd, } \frac{b^m - 1}{2} \geq n, \quad (8a)$$

$$\text{if } b \text{ is even, } \frac{b^m - 2^m}{2} \geq n. \quad (8b)$$

Rule 2. No characteristic may be repeated; i.e., each digit must have a characteristic different from that associated with any other digit.

Convention 1. The various digits of a characteristic are arranged in a set order; i.e., $C_{1i}, C_{2i}, \dots, C_{mi}$. The first digit which is neither zero, nor (in case b is even) $b/2$, must be less than $b/2$. There must be at least one such digit.

Convention 2. The characteristic of the j th check digit has a 1 in the j th position and 0's elsewhere.

Rule 1 is required since, for a code of this type, we must be prepared to correct any digit in one of two ways (± 1). This implies a minimum of $2n + 1$ values of the corrector, one for each possible correction, and one for the case of no corrections. This means that b^m , the number of possible

* The above distinction between rules and conventions will be observed throughout this paper.

values of the corrector, must be at least $2n + 1$, equation (8a). For even bases, we must reject all values of the corrector containing only the digits 0 and $b/2$ for representing error conditions for the following reasons: a positive error leads to a corrector that is the characteristic of the incorrectly received digit, and a negative error leads to the b -complement of such a characteristic. In order to have unique error correction, we must be able to distinguish between these two conditions. If a characteristic were to contain only the digits 0 and $b/2$, it would be equal to its own b -complement; such combinations of digits are therefore not useable as characteristics or characteristic complements.

Rule 2 is required to permit a unique identification of an incorrect digit in case of a single error.

Convention 1 allows us to distinguish between positive and negative errors. By observing this convention, a characteristic (corresponding to a positive error) can be distinguished from its complement (corresponding to a negative error) by inspecting the first digit of a corrector which is neither 0 nor $b/2$. A characteristic will have this digit less than $b/2$, a characteristic complement will have this digit greater than $b/2$. If the corrector is a characteristic, the correction is minus one; if it is a characteristic complement, it is plus one.

Once the characteristics have been chosen, the corresponding encoding procedure may be performed in the following manner: Let a_{ij} represent the j th digit of the characteristic of information digit x_i . Let z_j represent the check digit which has a characteristic containing a 1 in the j th position. If convention 2 has been observed, (9) can be used to calculate z_j :

$$\sum_{i=1}^{n-m} a_{ij}x_i = -z_j \bmod b. \quad (9)$$

An encoder calculates each z_j and inserts it into the message in those digit positions which have the characteristic of the j th check digit assigned to them.

In more general terms, we use implicit relations that are equivalent to the explicit equations given by (9). Letting x_i represent an information or a check digit, and letting C_{ij} represent the j th digit of the characteristic of the i th information or check digit, these formulas may be rewritten as

$$\sum_{i=1}^n C_{ij}x_i = 0 \bmod b. \quad (10)$$

At the receiver, the decoder calculates m different check sums. Let c_j

represent the check sum corresponding to the j th corrector term, and x_i' represent the received value of x_i : Then,

$$\sum_{i=1}^n C_{ij}x_i' = c_j \bmod b. \quad (11)$$

The difference between equations (10) and (11) is the result of any mutilations caused by the channel. If no error has occurred, all the c_j 's are 0; if an error of ± 1 has occurred, the m c_j 's will form the characteristic or the characteristic complement, respectively, of the incorrectly received digit.

One disadvantage of a systematic code is the discontinuity in the number of check states as a function of m , the number of check digits. For example, in decimal code one check digit is required for a message of up to four digits, and two check digits for up to forty-eight digits. Obviously, for a message of intermediate length, for example, twelve digits, many of the corrector states cannot be used for single error correction since they will not correspond to any single error. A more efficient code would be obtained if the check states were limited to a smaller number.

One method of reducing the number of check states is to perform the check in a different modulus than the modulus of the channel. In the single error detection code using a mixed digit, binary check information and quinary message information was conveyed by this digit. This code was more efficient than a systematic code because each message contained the minimum number of check states which is 2.

If a mixed digit, x , is composed of the two components (y, z) where y is the information state of the digit and z the check state, it is convenient to combine these two components to form x by means of the formula

$$x = \alpha y + z. \quad (12)$$

We calculate z by using a linear congruence equation modulo α .

The use of this formula permits a decoder to act on x' , the received value of x , directly, without first resolving x' into y' and z' , because (12) insures that $x' = y' \bmod \alpha$. This permits x' to be corrected directly and then resolved into its components.

As an example, consider a semi-systematic code for correcting a single small error in a decimal system, using a twelve digit message; ten of the digits are information digits and two are mixed digits, each conveying binary message information and quinary check information. (One of these binary digits might represent the sign of the number.)

With two quinary checks, twenty-five different check states are possible; for correcting single small errors in a twelve digit message, twenty-

TABLE IV — CHARACTERISTICS AND CHARACTERISTIC COMPLEMENTS,
SEMI-SYSTEMATIC DECIMAL CODE

Digit	Characteristics	Characteristic Complements
x ₁ (mixed digit)	1 0	4 0
x ₂ (mixed digit)	0 1	0 4
x ₃	0 2	0 3
x ₄	1 1	4 4
x ₅	1 2	4 3
x ₆	1 3	4 2
x ₇	1 4	4 1
x ₈	2 0	3 0
x ₉	2 1	3 4
x ₁₀	2 2	3 3
x ₁₁	2 3	3 2
x ₁₂	2 4	3 1

five corrector states are required, one for each of the two possible corrections (± 1) for each digit, and one for the case of a correctly received message. Characteristics may be chosen for the various digits in accordance with the rules and conventions outlined above in this case, since the check modulus is the same for both check digits. Consequently, it is no accident that these characteristics, shown in Table IV, are the same as those shown in Table III.

Let C_{i1} and C_{i2} represent the characteristic of the i th digit, and let y_1 and y_2 represent the two binary information digits. Then:

$$\sum_{i=3}^{12} C_{i1}x_i = -z_1 \bmod 5, \quad x_1 = z_1 + 5y, \tag{13}$$

$$\sum_{i=3}^{12} C_{i2}x_i = -z_2 \bmod 5, \quad x_2 = z_2 + 5y_2. \tag{14}$$

Because $x_1 = z_1 \bmod 5$ and $x_2 = z_2 \bmod 5$, these relations can be re-written implicitly to resemble equation (10):

$$\sum_{i=1}^{12} C_{i1}x_i = 0 \bmod 5, \tag{15}$$

$$\sum_{i=1}^{12} C_{i2}x_i = 0 \bmod 5. \tag{16}$$

At the decoder, the corrector c_1c_2 is calculated by:

$$\sum_{i=1}^{12} C_{i1}x_i' = c_1 \bmod 5, \tag{17}$$

$$\sum_{i=1}^{12} C_{i2}x_i' = c_2 \bmod 5. \tag{18}$$

If the corrector is 00, the message has been correctly received; otherwise, the corrector is either the characteristic or characteristic complement of the incorrect digit, from which plus one or minus one respectively must be subtracted as a correction.

Consider the general case. Let $x_1, x_2, \dots, x_k$ represent the k information digits; $y_1, y_2, \dots, y_m$ represent the information state of the m mixed digits, and $z_1, z_2, \dots, z_m$ represent the check state of the m mixed digits. In addition, let $\alpha_1, \alpha_2, \dots, \alpha_m$ represent the number base of $z_1, z_2, \dots, z_m$ respectively; $\beta_1, \beta_2, \dots, \beta_m$ represent the number of possible states of $y_1, y_2, \dots, y_m$ respectively, and $x_{k+1}, x_{k+2}, \dots, x_{k+m}$ represent the values of the mixed digits after the message has been encoded. (Note that for simplicity, a check digit is considered as a special case of a mixed digit; its information state is permanently 0.) The following encoding procedure may be used in which $x_1, x_2, \dots, x_k$ are used directly as part of the transmitted message. This is a semi-systematic code, which means that information digits are not changed in coding. To derive the mixed digits, the following formulas are used:

$$a_{11}x_1 + \dots + a_{1k}x_k = -z_1 \bmod \alpha_1 \quad (19-1)$$

$$x_{(k+1)} = y_1\alpha_1 + z_1 \quad (20-1)$$

$$a_{12}x_1 + \dots + a_{2k}x_k + a_{2(k+1)}x_{(k+1)} = -z_2 \bmod \alpha_2 \quad (19-2)$$

$$x_{(k+2)} = y_2\alpha_2 + z_2 \quad (20-2)$$

$$\vdots$$

$$\vdots$$

$$a_{j1}x_1 + \dots + a_{jk}x_k + \dots + a_{j(k+j-1)}x_{(k+j-1)} = -z_j \bmod \alpha_j \quad (19-j)$$

$$z_{(k+j)} = y_j\alpha_j + z_j \quad (20-j)$$

$$\vdots$$

$$\vdots$$

$$a_{m1}x_1 + \dots + a_{mk}x_k + \dots + a_{m(k+m-1)}x_{(k+m-1)} = -z_m \bmod \alpha_m \quad (19-m)$$

$$x_{(k+m)} = y_m\alpha_m + z_m. \quad (20-m)$$

In each case, the value of the check component z_j , of a mixed digit $x_{(k+j)}$ is determined by a formula involving the information digits and previously calculated mixed digits. Immediately after z_j has been determined, $x_{(k+j)}$ is calculated for possible use in calculating $z_{(j+1)}$. After the message has been completely encoded the following equations, analogous to (10), will be satisfied.

Let C_{ij} represent a_{ji} in equation (19-j). Then,

$$\sum_{i=1}^{k+m} C_{ij}x_i = 0 \bmod \alpha_j. \quad (21)$$

(Since $x_{(k+j)} = z_j \bmod \alpha_j$, substitution of $x_{(k+j)}$ for z_j in equation (19-j) will continue to satisfy the equation.)

At the decoder, equation (21) is changed to

$$\sum_{i=1}^{k+m} C_{ij}x_i' = c_j \bmod \alpha_j \quad (22)$$

In (22), x_i' represents the received value of x_i , and c_j represents the j th digit of the corrector. If all the digits have been correctly received, i.e., $x_i' = x_i$ for all values of i , then $c_1 = c_2 = \dots = c_m = 0$; [see equation (21)]. If x_h had been received incorrectly so that $x_h' = x_h + 1$, but all other digits had been correctly received, then the value of c_j (the j th digit of the corrector) would be calculated in the following manner:

$$\begin{aligned} c_j \bmod \alpha_j &= \sum_{i=1}^{k+m} C_{ij}x_i' \\ c_j \bmod \alpha_j &= \sum_{i=1}^{k+m} C_{ij}x_i + C_{hj} = C_{hj} \end{aligned} \quad (23)$$

Equation (23) proves that C_{hj} is actually the j th digit of the characteristic of x_h , because by definition, the characteristic of x_h is the value of the corrector when $x_h' = x_h + 1$, and all other digits have been correctly received. This means that the general term, C_{ij} of (21), is actually the j th digit of the characteristic of the i th digit and that this is a simple characteristic code.

For the case that $x_h' = x_h - 1$, the value of the corrector is such that if it were incremented, digit by digit, by the characteristic of x_h , the corrector would be composed only of zeros. Incrementing the corrector by the characteristic of x_h is equivalent to recalculating the corrector with x_h' increased by one, which in this case would amount to calculating the corrector for the case of a correctly received message. The latter is composed of all zeros [see (21)]. Thus, for the case of a single error of -1 , the corrector is the characteristic complement of the digit which is incorrectly received. For a semi-systematic or systematic code, the characteristic complement is an m digit word whose j th digit is the complement modulo a_j of the j th digit of the characteristic.

Equation (20-j) shows that generally $\alpha_j\beta_j$ cannot exceed b . (An exception is given below.) The maximum value of y_j is $\beta_j - 1$ since y is a

digit in the number base β_j . The maximum value of z_j is usually $\alpha_j - 1$, since z_j is a digit in the number base α_j . Thus,

$$x_{k+j} = y_j \alpha_j + z_j \leq b - 1, \quad (24)$$

$$(\beta_j - 1) \alpha_j + \alpha_j - 1 \leq b - 1, \quad (25)$$

$$\alpha_j \beta_j \leq b. \quad (26)$$

Equation (24) restates (19-j), and also states that the maximum value of any digit x , is $b - 1$, where b is the number base of the channel. In (25), the maximum values of y_j and z_j are substituted to yield the result shown in (26).

It was stated above that the maximum value of z_j is usually $\alpha_j - 1$. An exception occurs only in case z_j checks only itself and other mixed digits, the latter being restricted to fewer than $b - 1$ states. Under such circumstances, the value of z is sometimes restricted, so that even though z is calculated to satisfy a check, modulo α_j [see equation (19-j)], it cannot assume $\alpha_j - 1$ values. For example, a code for transmitting a single digit message over a decimal channel and permitting the correction of small errors, might use as the set of transmitted messages the digits, 0, 3, 6, 9. In this case, $\alpha = 3$ (any correct message satisfies the check $x = 0 \bmod 3$) and $\beta = 4$ since four different messages may be transmitted. In this case, z is restricted to the value 0 because the mixed digit checks only itself.

In order to correct single errors of ± 1 , using a simple characteristic code, it is necessary and sufficient that every characteristic be different from every other characteristic, and that it also be different from the complement of every other characteristic.

The following rules and conventions may be used to derive a set of characteristics which meet the requirements for a simple characteristic semi-systematic or systematic code for correcting small errors for any base $b \geq 3$ and an arbitrary length message. No set of conventions can be found which will lead to a more efficient code of this class, since the rules, not the conventions limit the efficiency of the code.

Rule 1. For an n digit message, including mixed digits, containing m mixed or check digits of which m_1 are associated with an even modulus, α , the inequality

$$(\alpha_1 \cdot \alpha_2 \cdot \dots \cdot \alpha_{m_1} - 2^{m_1})/2 \geq n \quad (27)$$

must be satisfied.

Rule 2. No characteristic may be repeated, i.e., each digit must have a characteristic different from that associated with any other digit.

Rule 3. Since the m th check is the last one to be calculated, and the

characteristic of the m th mixed digit must therefore contain only a single digit which is not 0, α_m must be greater than 2.

Convention 1. The various digits of a characteristic are arranged in a set order, i.e., $C_{i1}, C_{i2}, \dots, C_{im}$. The first digit which is neither 0 nor $\alpha_j/2$ must be less than $\alpha_j/2$. There must be at least one such digit.

Convention 2. The characteristic of the j th mixed digit has a 1 in the j th position and 0's elsewhere, provided that $\alpha_j \neq 2$. If $\alpha_j = 2$, the characteristic of this mixed digit has a 1 in the j th and m th positions, and 0's elsewhere.

Rule 1 is required because the number of possible corrector states is $\alpha_1 \cdot \alpha_2 \cdot \dots \cdot \alpha_m$, of which only those containing at least one digit which is neither 0 nor $\alpha/2$ can be associated with the $2n$ possible errors. The same reasons used for Rule 1 for the systematic code case are equally applicable here; a characteristic containing only the digits 0 or $\alpha_j/2$ in the j th position is not distinguishable from its complement.

Rule 2 is required to permit a unique identification of an incorrect digit.

Rule 3 is necessary to derive the sign of an error on the m th mixed digit.

The reasons for using Conventions 1 and 2 in the case of the systematic code are equally applicable in this case. For the case $\alpha = 2$, however, a special convention must be used to avoid a conflict with Convention 1.

The procedure for converting a set of characteristics into an error correcting code system is the same for a semi-systematic code as for a systematic code except that the following additional functions must be performed: the encoder must combine check states with information states to derive mixed digits, and the decoder must resolve mixed digits into information and check digits *after* it has performed its corrections.

By using these rules and conventions, the most efficient simple characteristic code can be determined. For messages of length n (including mixed or check digits), the following relations must be satisfied:

Let

$$P = \alpha_1 \cdot \alpha_2 \cdot \dots \cdot \alpha_m,$$

$$Q = \beta_1 \cdot \beta_2 \cdot \dots \cdot \beta_m,$$

$$m_1 = \text{number of even } \alpha\text{'s}.$$

Then:

$$(P - 2^{m_1})/2 \geq n, \tag{28}$$

$$\alpha_i \beta_i \leq b. \tag{29}^*$$

* For exceptions, see above.

TABLE V — DECIMAL ERROR CORRECTION CODES

n	P	$\frac{10^m}{Q}$	$\alpha_1, \alpha_2, \dots$	$\beta_1, \beta_2, \dots$	$2n + 1$
1	3	2.5	3	4	3*
2	5	5	5	2	5
3	10	10	10	1	7
4	10	10	10	1	9
5	15	16.7	5, 3	2, 3	11
6	15	16.7	5, 3	2, 3	13
7	15	16.7	5, 3	2, 3	15
8	20	20	10, 2	1, 5	17
9	25	25	5, 5	2, 2	19
10	25	25	5, 5	2, 2	21
11	25	25	5, 5	2, 2	23
12	25	25	5, 5	2, 2	25
13	30	33.3	10, 3	1, 3	27
14	30	33.3	10, 3	1, 3	29
15	40	40	10, 2, 2	1, 5, 5	31
16	40	40	10, 2, 2	1, 5, 5	33
17	50	50	10, 5	1, 2	35
18	50	50	10, 5	1, 2	37
19	50	50	10, 5	1, 2	39
20	50	50	10, 5	1, 2	41

* The single digit message containing the points 0, 3, 6, 9 is an exception to the inequality $\alpha\beta \leq b$, because the mixed digit checks only itself.

For the most efficient code b^m/Q should be minimized. This term represents the ratio of the number of possible messages for an n digit message with and without error correction. This is normally at least as great as $2n + 1$, the number of possible corrections on such a message.

Table V shows the most efficient decimal codes of this type for an n digit message, for values of n from 1 to 20. Where two or more different codes are equally efficient, the code with the fewest mixed digits is shown. It is easy to convert from a code using two mixed digits with $\alpha_1 = 5$, $\alpha_2 = 2$, to one using a check digit with $\alpha = 10$, or to make the inverse conversion, and to show that both codes are equally efficient.

IV. SINGLE ERROR CORRECTION CODES, UNRESTRICTED ERROR

The problem of correcting an unrestricted error on one digit of a message must be divided into two categories, depending on whether b is a prime number or a composite number. As will be seen, the error correction problem for prime bases is considerably simpler than that for composite bases. The method for correcting errors in prime number systems was discovered by Golay,⁶ although this did not come to the author's attention until after he had worked out the same method. The

adaptation to non-prime channel bases is believed to be novel. Since the adaptation makes use of the code for prime bases, both will be described.

4.1 Prime Number Base, Single Unrestricted Error Correction Code

This code depends upon a fundamental property of prime numbers, well known in number theory.⁹ Let p represent a prime number and d , c , and w represent non-negative integers less than p , related by the expression:

$$dw = c \bmod p. \quad (30)$$

If $d \neq 0$, then d and c uniquely determine w .

In order to have a simple characteristic systematic code for correcting unrestricted errors, it is necessary and sufficient that the set of characteristics shall have the property that all multiples of all characteristics are distinct. Equation (30) implies a unique correspondence between multiples of a characteristic and the characteristic itself, if we consider c to be the multiple, d the multiplying factor and w a digit of the characteristic. An error, d , is simply identifiable if a known digit of a characteristic is always 1. If each characteristic is distinct from every other and if a sufficient number of check digits are available, a simple characteristic code can be obtained. In the following set of rules and conventions which may be used for deriving a set of characteristics for a simple characteristic systematic code for correcting single unrestricted errors, p represents the prime number base of the channel. The number base of the channel must be prime, and the length of the message is arbitrary. Since the rules and not the conventions limit the efficiency of the code, no other set of conventions may be found which will lead to a more efficient code of this class.

Rule 1. For an n digit message, m check digits are required and m must satisfy the inequality

$$n \leq \frac{p^m - 1}{p - 1}. \quad (31)$$

Rule 2. Each digit must have a different characteristic.

Convention 1. The digits of a characteristic are arranged in a set order, i.e., $C_{i1}C_{i2} \cdots C_{im}$. The first digit which is not 0 must be 1.

Convention 2. The characteristic of the j th check digit has a 1 in the j th position and 0's elsewhere.

Rule 1 is required for a code for correcting single unrestricted errors since any digit must be correctable in one of $p - 1$ ways. This implies a minimum of $n(p - 1) + 1$ states for the corrector, one for each cor-

rection and one for the correct message. When m check digits are used, p^m corrector states are obtained.

Rule 2 and Convention 2 are the same for the single small error correction systematic codes. The same reasons apply for both cases.

Convention 1 is changed from the equivalent convention for the small error correction code, because the magnitude of the error, not only its sign, must be derivable for a code for correcting single unrestricted errors.

An encoder first encodes the message according to (32), where C_{ij} represents the j th digit of the characteristic of x_i ,

$$\sum C_{ij}x_i = 0 \bmod b. \quad (32)$$

The decoder calculates the corrector using the following formula where x_i' represents the received value of x_i ;

$$\sum C_{ij}x_i' = c_j \bmod b. \quad (33)$$

The decoder then examines the digits of the corrector in order. The first digit which is not 0 shows the magnitude, d , of the error. All digits are then divided by d (provided $d \neq 0$). (That division is unique, as shown by (30).) The result of this division is the characteristic of the incorrect digit, which is then corrected by subtracting d .

Consider a code for correcting a single unrestricted error in a six digit message for a base 5 channel:

$$6 \leq \frac{5^m - 1}{4}. \quad (34)$$

A value of 2 for m will satisfy equation (34). The characteristics are 14, 13, 12, 11, 10 and 01, the last two being check digit characteristics, for x_1 , x_2 , x_3 , x_4 , x_5 , and x_6 respectively. Here, x_1 , x_2 , x_3 , and x_4 are information digits. The encoding formulas are:

$$x_1 + x_2 + x_3 + x_4 = -x_5 \bmod 5, \quad (35)$$

$$4x_1 + 3x_2 + 2x_3 + x_4 = -x_6 \bmod 5. \quad (36)$$

The decoding and correcting formulas are: (x_i' is the received value of x_i)

$$x_1' + x_2' + x_3' + x_4' + x_5' = c_1 \bmod 5, \quad (37)$$

$$4x_1' + 3x_2' + 2x_3' + x_4' + x_6' = c_2 \bmod 5. \quad (38)$$

The corrector is c_1c_2 .

Suppose that a message 221321 is received as 224321. Then:

$$c_1 = 13 = 3 \bmod 5, \quad (39)$$

$$c_2 = 26 = 1 \bmod 5. \quad (40)$$

To find the characteristic of the digit, x_h , that was incorrectly received from the value of the corrector, (41) and (42) must be solved:

$$d C_{h1} = c_1 = 3 \bmod 5, \quad (41)$$

$$d C_{h2} = c_2 = 1 \bmod 5. \quad (42)$$

Because the first non-zero digit of any characteristic is 1, (41) can be solved for d since $C_{h1} = 1$. This yields the result, $d = 3$. Using this result, (42) is solved for C_{h2} ; by inspection, $C_{h2} = 2$, since $3 \cdot 2 = 6 = 1 \bmod 5$. Thus the characteristic of the incorrect digit, $C_{h1} C_{h2}$, is 12, and the error d , is 3; x_3' must therefore be reduced by 3 to get the correct value. Since the message was received with x_3' too high by an amount 3, this result confirms our expected correction.

Any correction that is applied must be applied on a modulo b basis. For example, if a correction of -2 is indicated on a digit whose received value is 1, $1 - 2 = 4 \bmod 5$, which means that the digit is corrected to 4.

Codes of this type are restricted in their construction. No mixed digits may be used, and the number base must be prime. For the case of $n = [(p^g - 1)/(p - 1)] + 1$, $g + 1$ check digits are required [see (31)]. This means that the number of information digits for a message of this length is the same as for a message one digit shorter, which requires only g check digits. A comparable binary case is the Hamming Code example of an eight binary digit message (four information digits) compared with a seven digit message (also four information digits). In the binary case, the extra digit is useful for double error detection, but unfortunately, this is not the case for non-binary codes.

4.2 Composite Number Base, Single Unrestricted Error Correcting Code

The problem of correcting an unrestricted error on a single digit, working with a number base b , that is not a prime is much more difficult. Many relatively inefficient techniques exist. For example, characteristics containing only binary numbers (0 and 1) might be used; (this would amount to using the Hamming Code directly). This is obviously inefficient since the corrector associated with any single digit error of amount

d , would contain only the digits 0 and d , thus wasting most of the possible corrector values.*

It is possible to encode and decode using the prime factors of the number base, performing separate and independent corrections on each factor. This is also inefficient, since for many cases, information as to which digit is in error is found independently in two or more ways, while for certain values of the error, it can be found in only one way. Working with mixed digits and check bases, α lower than b , is not satisfactory since certain values of the error (α in particular) will never show up in a particular check. The technique used for primes will not work since multiples of two different characteristics may be identical; for example, base 10, characteristics 11 and 13, error 5, will both yield correctors of 55.

Another technique that is relatively efficient is, however, available. It involves performing all check, encoding and decoding operations in a number base p , where p is some prime number (usually, the lowest) that is equal to or greater than b . (In case b is a prime, we use the procedure outlined above, which is a special case of the procedure to be described below.)

The obvious difficulty in such a procedure is that while the information channel can only handle b levels, the check digits may assume p levels, corresponding to the required p check states. This dilemma can be resolved by adding an *adjustment digit*. The object of this digit is to permit check information to be transmitted in a base greater than b , the channel base. The idea of an adjustment digit can best be illustrated by an example. Suppose for a decimal channel, checks are performed in a unodecimal (base 11) code. Let γ represent the value corresponding to ten. (The consecutive integers in a unodecimal system are then 0, 1, 2, 3, $\dots$, 9, γ , 10, 11, $\dots$, 19, 1γ , 20, etc.) Suppose in a particular message, four check digits, z_1, z_2, z_3, z_4 , calculated modulo 11 from decimal information digits are used, whose values are 1, 0, γ , 8. A fifth digit, z_0 is added such that the sums modulo 11 of $z_1 + z_0, z_2 + z_0, z_3 + z_0, z_4 + z_0$ are kept constant at 1, 0, γ , 8 respectively. There are eleven different words satisfying the condition: $[1, 0, \gamma, 8] = [(z_1 + z_0), (z_2 + z_0), (z_3 + z_0), (z_4 + z_0)]$. These are shown in Table VI. Of these words, six do not contain the digit γ , and so may be transmitted over a decimal channel. Thus, an adjustment digit permits check digits which are calculated in a number system of a higher base than b , to be transmitted over a base b channel. When an adjustment digit is used in base p for adjusting m digits so that transmission over a channel in base b is possible, a mini-

* A waste of corrector values is equivalent to an excessive number of check states for a message, which in turn implies an excessive number of check digits.

mum of $b - m(p - b)$ states are allowed for the adjustment digit. (For certain values of the check digits, more states could be allowed, but a code for utilizing these extra states becomes unwieldy.) For the case $b = 10$, $p = 11$, this turns out to be $10 - m$. At least one state must be available for each adjustment digit, to have a workable code.

The characteristic of an adjustment digit is determined in the following way: if an adjustment digit adjusts the j th check digit, then the j th digit of the characteristic of the adjustment digit is 1; otherwise, it is 0. The characteristic of all other digits may be derived using the rules described above for the prime number base channel, except that p , the prime number base of the code must be used instead of b , the number

TABLE VI — ILLUSTRATION OF ADJUSTMENT DIGIT

z_0	z_1	z_2	z_3	z_4
0	1	0	γ	8
1	0	γ	9	7
2	γ	9	8	6
3	9	8	7	5
4	8	7	6	4
5	7	6	5	3
6	6	5	4	2
7	5	4	3	1
8	4	3	2	0
9	3	2	1	γ
γ	2	1	0	9

base of the channel, for generating characteristics. A message is initially encoded using a value of 0 for an adjustment digit. Subsequently, if the adjustment digit always has at least q allowable states, it may be used to transmit one additional information digit, base q , of information. If the value of this information digit is y , the $(y + 1)$ st lowest possible value of the adjustment digit (making the lowest value equivalent to $y = 0$) meeting the requirement that all adjusted check digits are no greater than $b - 1$ is transmitted. The adjustment digit in conjunction with its associated check digits conveys a digit, base q , of information.

In the example given above, $q = 6$ and if y is 4, the fifth lowest value of z_0 , 7, is transmitted. The lowest value must be associated with $y = 0$. The values of $z_0z_1z_2z_3z_4$ that are sent over the decimal channel are 75431.

An example of such a code is one using a decimal channel working in a unodecimal base for the purposes of encoding and error correction. The word length, n , is twelve, nine decimal information digits, one octal (base

8) information digit associated with the adjustment digit, and two check digits. The characteristics are the following:

x_1 1 γ	x_5 16	x_9 12
x_2 19	x_6 15	x_{10} 11 (adjustment digit)
x_3 18	x_7 14	x_{11} 10 (check digit)
x_4 17	x_8 13	x_{12} 01 (check digit)

Let z_{11} and z_{12} represent the values of the check digits x_{11} and x_{12} , originally derived from $x_1, x_2, \dots, x_8, x_9$:

$$x_1 + x_2 + x_3 + \dots + x_9 = -z_{11} \bmod 11, \tag{43}$$

$$\gamma x_1 + 9x_2 + 8x_3 + \dots + 2x_9 = -z_{12} \bmod 11. \tag{44}$$

From z_{11} and z_{12} , the ten different words $(0, z_{11}, z_{12}), (1, z_{11} - 1, z_{12} - 1), (2, z_{11} - 2, z_{12} - 2), \dots, (9, z_{11} - 9, z_{12} - 9)$ are formed. If y is the value of the octal information digit, the $(y + 1)$ st such word, that does not contain the digit γ , is selected and transmitted as the last three digits of the message. For example, if $z_{11} = 2, z_{12} = 1$ and $y = 6$, the ten words are $(0, 2, 1), (1, 1, 0), (2, 0, \gamma), (3, \gamma, 9), (4, 9, 8), (5, 8, 7), (6, 7, 6), (7, 6, 5), (8, 5, 4), (9, 4, 3)$; the word $(8, 5, 4)$ is selected since it is the seventh in the sequence that does not contain any γ 's. Table VII shows the choice of the three last digits as a function of y , given $z_{11} = 2, z_{12} = 1$.

Formula (45) is used for calculating the corrector. Let C_{ij} represent the j th digit of the characteristic of x_i , c_j the j th digit of the corrector, and x_i' the received value of x_i . Then,

$$c_j = \sum_{i=1}^{12} C_{ij}x_i' \bmod 11. \tag{45}$$

The translation from corrector to correction is the same as if the original

TABLE VII — RELATION BETWEEN ADJUSTED DIGIT AND ASSOCIATED INFORMATION

y	x_{10}	x_{11}	x_{12}
0	0	2	1
1	1	1	0
2	4	9	8
3	5	8	7
4	6	7	6
5	7	6	5
6	8	5	4
7	9	4	3

message had been in a unodecimal code. (This has been illustrated in Section 4.1.)

The first step of the encoding procedure is to calculate the unadjusted check digits. Next, the adjusted check digits and adjustment digit are selected according to the value of y , the information digit associated with the adjustment. The message is then ready for transmission.

At the decoder, the message is first corrected as if it had been received as a unodecimal message. The information digits are then in their corrected states. Next, the adjustment digit and the check digits are examined and the inverse of the encoding process used to select a particular set of check and adjustment digits is used to reconstruct the value of y which originally controlled the selection. In the example given above, the values of x_{10} , x_{11} , x_{12} are 8, 5, 4 respectively; the decoder recognizes that this is the seventh lowest value of x_{10} , which means that the value of y , used in selecting x_{10} and the adjusted values of x_{11} and x_{12} , was 6.

The code described above is fairly efficient; about 90 per cent of the corrector values can be associated with corrections; the product of the information states and the check states is about 97 per cent of the total number of states of a twelve decimal digit word. Each of the above factors reduces the efficiency of the code below a possibly unattainable maximum. It will be noted, however, that this reduction is relatively small in both cases, and is very much lower than would be the case for any of the rejected schemes. The scheme is not difficult to instrument; relatively little additional equipment is required in addition to the basic equipment for instrumenting a simple prime number base channel, unrestricted single error correcting code system.

The method of adjustment digits is general and can be used for deriving a single error correction code for correcting unrestricted errors for any channel base. Any convenient prime check base, p , at least as great as b may be used, although the lowest will generally be the most efficient. The only requirements which must be fulfilled are that the number of states of the adjustment digit must be at least 1, and that at least two check digits must be associated with each adjustment digit. An adjustment digit associated with m check digits, working with a channel base b and a check base p , may have $b - m(p - b)$ different states.

V. SINGLE ERROR CORRECTION, DOUBLE ERROR DETECTION CODES FOR CORRECTING SMALL ERRORS

Single error correction, double error detection codes are very useful in situations where a message may occasionally be repeated. In order

for a correction code to be reasonably useful in a system with random noise or errors, the errors must be relatively infrequent, which makes double errors still more infrequent. If means are available for an occasional but very infrequent repetition of a message, a single error correction, double error detection code will increase the reliability of a digital system, since a message may be repeated if a double error is recognized.

This section will show how the ideas of the single error correction, double error detection Hamming Code may be combined with the ideas of semi-systematic single small error correction codes (described in Section III) to derive simple and efficient codes for correcting single small errors and detecting double small errors.

In order to derive a simple characteristic code for correcting single small errors, and detecting double small errors, a set of characteristics must be found having the property that the sum or difference of two characteristics or their complements or double the value of one characteristic or its complement be distinguishable from the value of any single characteristic or its complement. The sum of two characteristics represents the value of the corrector for a message with two errors of $+1$, $+1$, the difference represents two errors of $+1$, -1 , the sum of their complements represents two errors of -1 , -1 ; double a characteristic represents an error of $+2$, and double a complement represents an error of -2 . To have a true single error correction, double error detection code for small errors, all these cases must be distinguished from the case of a single error or no error by making certain that the value of the corrector for any of these cases is different than the value of the corrector corresponding to any single error and no error.

Table VIII gives the characteristics used in the single error correction Hamming Code and the single error correction, double error detection Hamming Code for conveying four digits of information in a message containing seven or eight binary digits respectively.

An inspection of Table VIII shows that the sum (performed without carries from column to column) of any two characteristics in the right part of the table is distinguished by having at least one 1 in the first three places and a 0 in the last place. This distinguishes it from any single characteristic since all characteristics have a 1 in their last place.

Some difficulties arise in trying to adopt such a scheme directly in a non-binary system. For the code to be efficient, an over-all check would have to be performed using a mixed digit; only two check states are required for an over-all parity check, and if $b > 3$, (b representing the number base of the channel) at least two information states are possible. But the over-all check digit, which performs a binary check, is not checked by any other digit. This means that although errors might be

detected in an over-all check digit, difficulties would be encountered in determining the direction of the correction, so that the information conveyed by the mixed digit could be used. Actually, means are available, for accomplishing an adaptation of binary techniques. These methods are described in Section VII but they are less straightforward than the ones described below.

For channels with base b , greater than 3, at least one check may be made using a check base, α_m , that is 4 or greater. If characteristics are used whose last digit (the digit associated with the α_m check) is always 1, and whose only other limitation is that each characteristic is different from every other characteristic, a satisfactory code is obtained. Single errors are corrected in the normal way. If the last digit of the corrector is 1 or $\alpha_m - 1$, the error is ± 1 respectively on the digit whose charac-

TABLE VIII — CHARACTERISTICS FOR HAMMING CODES

Single Error Correction			Single Error Correction Double Error Detection
001	Check Digit	x_1	0011
010	Check Digit	x_2	0101
011	Information Digit	x_3	0111
100	Check Digit	x_4	1001
101	Information Digit	x_5	1011
110	Information Digit	x_6	1101
111	Information Digit	x_7	1111
	Over-all Check Digit	x_8	0001

teristic or whose characteristic complement is indicated by the corrector. If the last digit of the corrector is 2 or $\alpha_m - 2$, or the last digit is 0 and other digits are not all 0, a double error is indicated. If the entire corrector is made up of 0's, the message is correct as received.

An example is a code for a ten digit message, decimal base channel; eight decimal information digits, one mixed digit conveying binary message information (such as the sign of the decimal number) and quaternary (base 4) check information, and one check digit are transmitted in each message. Let x_1 and x_2 represent the mixed and check digit respectively, x_3 through x_{10} the information digits, y_1 the binary information conveyed by x_1 , and z_1 the quaternary check information conveyed by x_1 . The encoding formulas are:

$2x_3 + 3x_4 + 4x_5 + 5x_6 + 6x_7 + 7x_8 + 8x_9 + 9x_{10} = -x_2 \bmod 10,$ (46)

$x_2 + x_3 + x_4 + x_5 + x_6 + x_7 + x_8 + x_9 + x_{10} = -z_1 \bmod 4,$ (47)

$x_1 = z_1 + 4y_1.$ (48)

Note that (46), (47) and (48) must be applied consecutively, in that order, since (47) cannot be applied without knowing x_2 obtained from (46), and (48) requires z_1 , obtained from (47).

The characteristics are 01, 11, 21, 31, 41, 51, 61, 71, 81, 91 respectively; the complements of the characteristics are 03, 93, 83, 73, 63, 53, 43, 33, 23, 13 respectively. The corrector, c_1c_2 , is calculated at the decoder by the following formulas (x_i' is the received value of x_i):

$$c_1 = \sum_{i=2}^{10} (x_i')(i-1) \bmod 10 \quad (49)$$

$$c_2 = \sum_{i=1}^{10} x_i' \bmod 4 \quad (50)$$

Consider the example of a message with decimal information digits 3 7 5 2 0 6 5 2 and binary information digit 1. Then $x_2 = 3$, $z_1 = 3$, and $y_1 = 1$, yielding a value of 7 for x_1 . The message is sent as 7 3 3 7 5 2 0 6 5 2. Suppose that the sixth digit is changed to 1 in transmission. Then the corrector has a value 53; this is the complement of the characteristic of the sixth digit and indicates that the sixth digit should be incremented by 1 according to the rules previously stated. If the sixth digit had been received as 1 and the seventh digit also received as 1 (an error of $+1$), then the corrector value would be 10, indicating a double error (see rules stated above).

If a multiple of 4 is used as α_m , the last digit of a characteristic may assume all odd values below $\alpha_m/2$. The rule then is that an even value of the last digit of the corrector, or a 0 for the last digit and other digits of the corrector not all 0, indicates a double error.

The following set of rules and conventions may be used with any base $b \geq 4$, and any length of message, for deriving a set of characteristics for a semi-systematic code for correcting single small errors and detecting double small errors. Since the conventions restrict the efficiency of the code, it is conceivable that a different set of conventions will yield a more efficient code in some cases; (51) may be modified through the use of an alternate set of conventions.

Rule 1. No two digits may have identical characteristics.

Convention 1. Choose for α_m a multiple of 4. Let $\alpha_m/4 = g$.

Convention 2. The characteristic of the mixed digit associated with α_m contains a single 1 in the last position; the rest of its digits are 0.

Convention 3. The characteristics of the j th mixed or check digit contains a 1 in the last position, a 1 in the j th position and 0's elsewhere.

Convention 4. The characteristic of an information digit has an odd

number less than $\alpha_m/2$ in its last position. The rest of its digits are arbitrary.

Convention 5. The above conventions restrict the choice of characteristics. In order to have n distinct characteristics, m mixed or check digits, using check bases $\alpha_1, \alpha_2, \dots, \alpha_m$, are required, and inequality (51) must be satisfied:

$$n \leq \alpha_1 \alpha_2 \dots \alpha_{m-1} \cdot g. \quad (51)$$

Codes may be derived using the above conventions only if $b \geq 4$. For the ternary case, a relatively efficient code may be obtained by using one ternary digit as an over-all parity check digit. The rest of the message is in a single small error correction code, derived using the rules and conventions of Section III. Any single small error will lead to a failure of the parity check, and a double small error will lead to a failure of other checks but not the parity check.

No general solution has been found for deriving an efficient single error correction double error detection code for the unrestricted error case. Also, no general solution has been found for deriving an efficient multiple error correction code for the unrestricted error case. A reasonably efficient method has been found for correcting multiple errors in the more important small error case; this is discussed in Section 6.2.

VI. THE USE OF BINARY ERROR CORRECTION TECHNIQUES IN NON-BINARY SYSTEMS

In this section, methods for using binary codes for the correction of errors in a non-binary system are described. Although the single small error correction codes obtained in this manner are generally less flexible than the codes obtained in Section III, the class of multiple error correction codes described in Section 6.2 is the only reasonably satisfactory class of such codes that has been found. The codes described in this section are semi-systematic but are not simple characteristic codes.

6.1 *Single Small Error Correction Codes*

Binary codes are most conveniently used for correcting small errors (± 1). Suppose any digit, base b , has an associated pair of binary digits, arranged in such a way that a change of ± 1 in the base b digit will change only one of the two binary digits. For $b = 10$, an association such as the one shown in Table IX might be used. For example, if a 6 is received as a 7, the associated binary message would indicate that the second of the binary digits is incorrect; a 7 can be corrected

TABLE IX — ASSOCIATED BINARY DIGITS FOR CORRECTION OF SMALL ERRORS

Decimal Digit	Associated Binary Digits
0	00
1	01
2	11
3	10
4	00
5	01
6	11
7	10
8	00
9	01

TABLE X — REFLECTED QUIBINARY CODE

Decimal Digit	Quinary Component	Binary Component	Associated Binary Digits
0	0	0	00
1	0	1	01
2	1	1	11
3	1	0	10
4	2	0	00
5	2	1	01
6	3	1	11
7	3	0	10
8	4	0	00
9	4	1	01

to an 8 or a 6, but only the correction to 6 would correspond to a change in the second binary digit of the associated binary message.

If the first of the associated binary digits is the odd or even indication of a quinary component of a decimal digit, a decimal digit can convey ten states rather than the four states of the associated binary digits. The combination of binary and quinary digits shown in Table X may be called a reflected quibinary code because of its analogy with the reflected binary code.*

If a method were available for transmitting without error (e.g., by using an error correcting code) a message composed of the associated binary digits in a base b code, small errors could be corrected in the base b digits.

An examination of Table X for resolving a decimal digit into binary and quinary components, reveals that a change of ± 1 on any decimal

* The reflected binary code has the property that each increment changes only one binary digit; for example, the eight successive words of a three binary digit reflected binary code are 000, 001, 011, 010, 110, 111, 101, 100.

digit will change only one of these two components. Further, an error corresponding to a change in the quinary component can be uniquely corrected if the error in the decimal digit is assumed to be ± 1 . For example, if a received 6 is discovered to have an incorrect quinary component, only a decrease in the quinary component making the decimal digit 5 is a possible correction, since an increase in the quinary component would correspond to the decimal digit 9, a change of more than ± 1 from 6.

A system is shown in Fig. 2 for taking advantage of these properties.

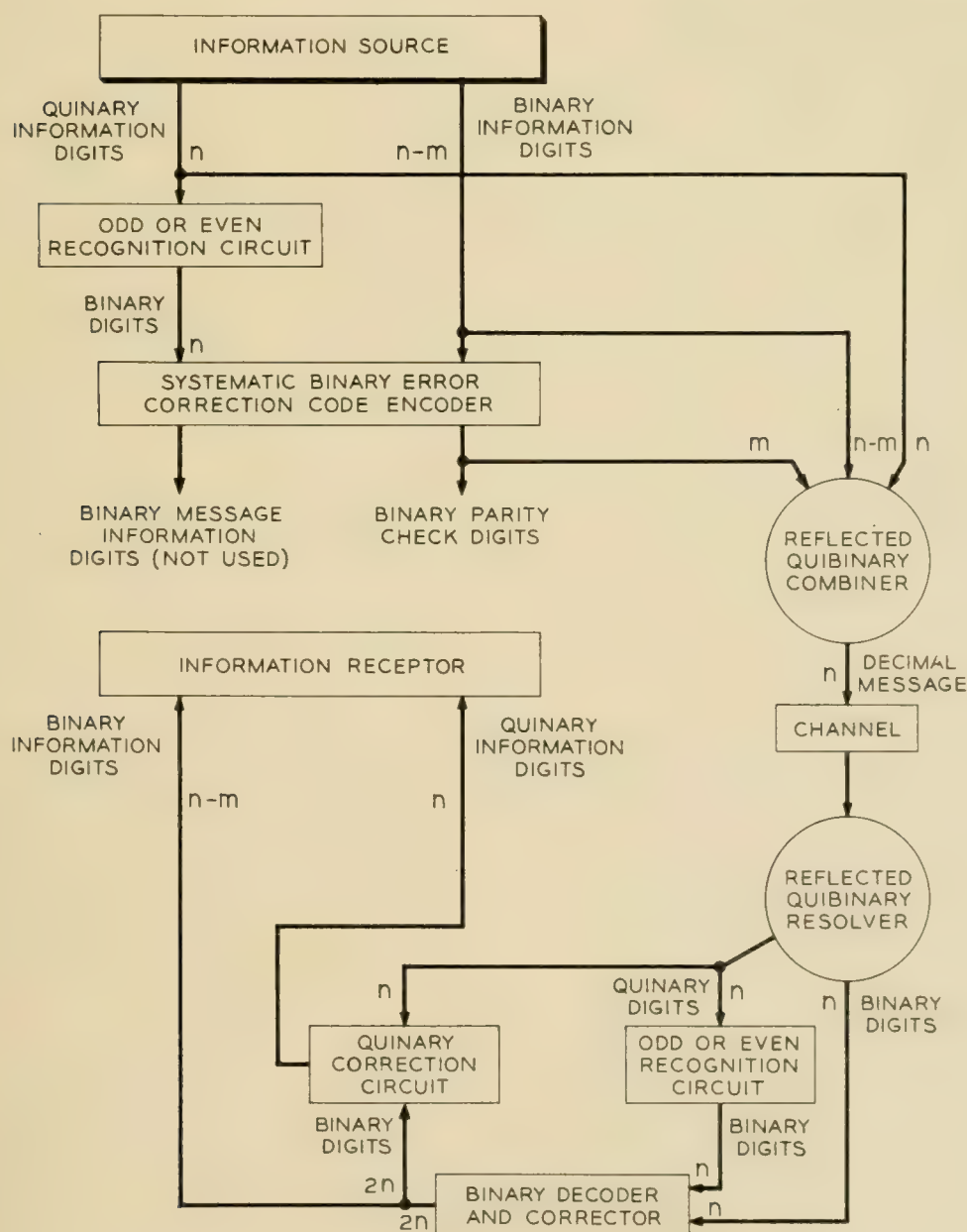


Fig. 2 — Use of binary codes with a decimal channel.

In this example, an information source generates n quinary and $n - m$ binary information digits for each message. All quinary digits go through an odd or even recognition circuit to be converted into binary digits for the purpose of generating a binary error correction code message. These binary digits and the binary digits generated by the information source are fed into a systematic binary error correction code encoder whose output is a binary message containing $2n$ digits, of which m are parity check digits. This output is divided into two parts, $2n - m$ original inputs to the encoder unchanged by the encoding process (this is a systematic encoder which does not change information digits in encoding), and m parity check digits.

The m parity check digits are then combined with m of the quinary information digits through the use of the reflected quibinary combiner to form m of the decimal digits of the decimal message that is transmitted; the other decimal digits are formed by combining the $n - m$ binary information digits with the rest of the quinary information digits.

The decimal message is transmitted over the noisy channel and arrives with one or more (a number limited by the choice of the binary code) errors of ± 1 on decimal digits. It is fed into a reflected quibinary resolver which resolves decimal digits into binary and quinary components in accordance with the reflected quibinary code (Table X). The quinary digits are then fed into an odd or even recognition circuit to form binary digits; these and the binary outputs of the resolver are fed into a binary decoder and corrector, working with the same code as the binary encoder. The output of this corrector should correspond to the output of the original binary encoder.

In the decoder, the binary digits are corrected. When the binary digit derived from a quinary digit is corrected, however, the quinary digit is not yet correct. The correction of the quinary digit is performed by examining both the corrected binary digit derived from the quinary digit and the corrected binary digit which was derived from the same decimal digit as the quinary digit in question. The rules for correcting the quinary digit are given in Table XI.

As an example, consider the application of a Hamming Code for transmitting ten binary digits in a fourteen binary digit message.

Using a code of this type, single errors of ± 1 may be corrected in a seven digit decimal message, transmitting seven quinary digits of information and three binary digits of information. The characteristics required for a fourteen binary digit Hamming Code message are shown in the first column of Table XII.

TABLE XI — CORRECTING QUINARY DIGITS

<i>Q</i>	<i>B</i> ₁	<i>B</i> ₂	Correction of Quinary Digit
Even	0	0	None
Even	0	1	None
Even	1	0	-1
Even	1	1	+1
Odd	0	0	+1
Odd	0	1	-1
Odd	1	0	None
Odd	1	1	None

TABLE XII — BINARY CODE USED FOR CORRECTING
DECIMAL MESSAGE

Binary Characteristics		<i>a</i>	<i>b</i>	Position in Decimal Message
0 0 0 1	Parity Check Digit	(0)	(0)	Binary comp. of 1st digit
0 0 1 0	Parity Check Digit	(0)	(0)	Binary comp. of 2nd digit
0 0 1 1		1	1	Binary comp. of 3rd digit
0 1 0 0	Parity Check Digit	(1)	(1)	Binary comp. of 4th digit
0 1 0 1		0	0	Binary comp. of 5th digit
0 1 1 0		0	0	Binary comp. of 6th digit
0 1 1 1		1	3	Quinary comp. of 7th digit
1 0 0 0	Parity Check Digit	(1)	(1)	Binary comp. of 7th digit
1 0 0 1		1	3	Quinary comp. of 6th digit
1 0 1 0		0	2	Quinary comp. of 5th digit
1 0 1 1		0	4	Quinary comp. of 4th digit
1 1 0 0		1	1	Quinary comp. of 3rd digit
1 1 0 1		1	3	Quinary comp. of 2nd digit
1 1 1 0		0	0	Quinary comp. of 1st digit

To illustrate the method completely, a strictly binary example will first be illustrated, then a related decimal example. In column *a* of Table XII, the digits of a binary message are indicated and in column *b*, the binary and quinary information digits. The values of the parity check digits, which are shown in parentheses, are calculated by the usual formula. Let *C_{ij}* represent the *j*th digit of the characteristic of the *i*th digit (including parity check digits):

$$\sum_{i=1}^{14} x_i C_{ij} = 0 \text{ mod } 2.$$

(52)

This formula applies for all values of *j* and in this case will yield four implicit equations each with one unknown term, the value of the parity check digit. Using the given values of the binary information digits, the values of the parity check digits are calculated. These are shown in parentheses in Table XII.

The binary message is

0 0 1 1 0 0 1 1 1 0 0 1 1 0.

For this example, the quinary components (quinary information digits) of decimal digits are chosen odd if the corresponding digit of the binary example is 1, even if that digit is 0. The binary and quinary components are then combined by the rules of the reflected quibinary code to form the decimal digits 0 7 2 9 4 7 6. For example, the quinary and binary components of the fifth digit are 2 and 0, respectively; the decimal digit which has these components is 4, the fifth decimal digit of the message.

Consider the binary case. Suppose that the message is mutilated in transmission so that the tenth digit is received incorrectly. The message is mutilated from

0 0 1 1 0 0 1 1 1 0 0 1 1 0

to

0 0 1 1 0 0 1 1 1 1 0 1 1 0.

The decoder and corrector calculates the corrector by

$$c_j = \sum_{i=1}^{14} x'_i C_{ij} \bmod 2. \quad (53)$$

In this formula, c_j is the j th digit of the corrector and x'_i the received value of x_i . In this example the corrector is 1 0 1 0, which means that the tenth digit, which has this characteristic, is wrong and should be changed to 0.

The corresponding error in the decimal example is a change in the fifth digit from 4 to 3. If the message 0 7 2 9 3 7 6 is received, the resolver and quinary to binary converter delivers the message

0 0 1 1 0 0 1 1 1 1 0 1 1 0

to the decoder instead of

0 0 1 1 0 0 1 1 1 0 0 1 1 0

corresponding to the correct message. The corrected binary message is produced at the output of the decoder and corrector. When the quinary and binary components of the fifth digit are examined by the quinary correction circuit, the following inputs exist:

Received quinary digit	1 (Odd) (quinary component of received decimal 3)
Corrected binary digit derived from quinary	0 (B_1)
Corrected binary digit from same decimal number	0 (B_2).

Table XI shows that the quinary digit must be increased by 1 to 2, which combined with the binary 0 conveyed by the same decimal digit yields a decimal value of 4, the original transmitted value.

The best semi-systematic simple characteristic code for correcting single small errors in a seven digit message allows 6×10^5 possible messages in a seven digit message (see Table V), whereas this code allows 6.25×10^5 . This code is therefore slightly more efficient. In addition, this code has the special advantage that any error of ± 2 on one digit is recognizable since the corrector will have a value of 1111 for the associated binary message. (An inspection of the choice of characteristics and assignment of characteristics to the two components of any decimal digit will confirm this.)

This general technique can be applied to any base b channel, provided

TABLE XIII — COMPONENTS OF QUINARY DIGITS

Mixed Digit			Information Digit		
Quinary Digit	Info. Comp.	Check Comp.	Quinary Digit	Binary Comp.	Ternary Comp.
0	0	0	0	0	0
1	0	1	1	1	0
2	1	1	2	1	1
3	1	0	3	0	1
4	not used	not used	4	0	2*

* If quinary information is initially generated, the combination (1, 2) will not occur.

that b is greater than 3. For odd bases, the digits which convey a parity check component and an information component cannot be utilized efficiently since one state of the base b digit is not available. For example, using a base 5, (see Table XIII), only two information and two parity check states may be conveyed by one digit, since the use of a third information state would require at least six states for the mixed digit. In the case of information digits, however, all states can be used. In the quinary example, the resolution of a digit into two components and the subsequent recombination is subject to the restraint that one of the combinations (1, 2) will not occur, which can be assured if the information source generates quinary digits.

For the case of high redundancy codes having the property that the associated binary code contains more than 50 per cent parity check digits (corresponding to a negative value of $n - m$ in Fig. 2), at least some of the base b digits must convey two or more parity check digits.

This can be easily accomplished: a decimal digit can convey three

TABLE XIV — DECIMAL DIGIT CONVEYING THREE BINARY DIGITS

Decimal Digit	Binary Components
0	0 0 0
1	0 0 1
2	0 1 1
3	0 1 0
4	1 1 0
5	1 1 1
6	1 0 1
7	1 0 0
8	not used
9	not used

parity check digits if a simple reflected binary code correspondence between binary and decimal digits is maintained as shown in Table XIV.

An extension of this idea is the encoding of the original information (i.e., the information that is shown coming out of the information source in Fig. 2) in some error detection or correction code. For example, the decimal to reflected quibinary code resolver will cause both components to be incorrect if an error of ± 2 in a decimal digit occurs. In this case, the system shown in Fig. 2 will automatically make a correction on the decimal digit of either $+2$ or -2 depending upon the value of the received decimal digit, and provided a double error correction binary code is used. Such a correction will be incorrect about half the time. If the received binary digit is compared to the corrected binary digit and the received quinary digit is compared to its corrected odd or even digit, an error of ± 2 can be detected without changing the code. If one extra binary check digit, treated as an information digit by the encoder and decoder, is transmitted in the message, this binary digit can convey the information necessary for determining the sign for a correction of ± 2 , provided that only one such correction is required for any one message. A rule for determining the value of this digit is:

$$\begin{aligned} B_c = 0 \quad & \text{if} \quad \sum_{i=1}^n q_i = (0 \text{ or } 1) \bmod 4, \\ B_c = 1 \quad & \text{if} \quad \sum_{i=1}^n q_i = (2 \text{ or } 3) \bmod 4, \end{aligned} \tag{54}$$

where q_i represents the i th quinary information digit, and B_c represents the special check digit. If the received message contains one error of ± 2 on a digit, two possible corrections may be made on the quinary component of this digit; ± 1 . Obviously, only one of these corrections will satisfy the equation for determining B_c since the two possible corrected values of q are two units apart.

Note that the associated binary codes for performing such a correction must have the property that two binary digits may be corrected since an error of ± 2 corresponds to incorrect values for two associated binary digits. If the noise is such that errors of ± 2 are not very unlikely, it may be desirable to place the binary and the quinary components of any one decimal digit in a different binary error correction code word so as to make the errors independent. In a seven decimal digit message, as an example, the quinary components of the first four decimal digits can be used to generate parity check digits which are conveyed by the binary components of the last three decimal digits. The binary component of the fourth decimal digit (this might be B_c) and the quinary components of the last three decimal digits generate parity check digits conveyed by the binary components of the first three decimal digits. Two separate binary error correction code messages are then conveyed by a single seven digit decimal code message. Each message is in a four information digit, three parity check digit Hamming Code. Through the use of this code, one error in the binary component of any decimal digit, and one error in the quinary component of any decimal digit may be corrected.

In certain cases, the quinary digits themselves might be encoded in an error correction code for single unrestricted errors before the binary process is carried out. This is helpful chiefly for occasional large errors, leading to initial miscorrections.

The variations based upon the principles described, which can be applied to any channel, provided $b \geq 4$, including the pyramiding of one code scheme upon another, are almost endless. Generally, the last encoding and first decoding step should be able to correct many more errors than the first encoding step. For example, if quinary components are encoded in single unrestricted error correction quinary code, the binary code should probably be a triple or quadruple error correction code; otherwise a correction may not correspond to the most probable error condition, and the correction scheme loses its effectiveness.

These techniques cannot be conveniently applied to the ternary channel, since a ternary digit cannot be resolved into two components efficiently.

6.2 *Multiple Small Error Correction Codes*

One limitation of the above techniques is the requirement for a systematic binary code; i.e., a code in which some of the binary information digits are transmitted directly, and others are determined by parity checks on information and previously calculated check digits. These

TABLE XV — REED-MULLER CODES — 256 DIGIT MESSAGE

Number of Digits of Information per Message	Number of Errors Correctable per Message
255	0
247	1
219	3
163	7
93	15
37	31
9	63
1	127

systematic codes are conveniently applicable only to the correction of single errors and a few special cases of multiple errors.

The Reed-Muller¹⁰ codes are not systematic codes, (“systematic” being used in the narrow sense indicated above, not in the sense of Hamming¹¹), but offer the advantage that multiple error correction is relatively straightforward. For this reason, it is desirable to find some way of adapting the binary Reed-Muller codes for correcting a number of small errors in non-binary codes.

To explain the nature of the Reed-Muller codes completely is beyond the scope of this paper; a list of their important features is sufficient. This is:

1. The length of a message is 2^k binary digits for the simpler versions of the code.

2. If C_c^d represents the number of combinations of d items taken c at a time, and $C_c^d = d!/[c!(d - c)!]$, then $2^k - \sum_{i=0}^m C_{k-i}^k$ information digits may be transmitted correctly in a message containing 2^k digits, if no more than $2^m - 1$ errors occur in the messages; 2^m errors are detected but they are not always correctable. The Reed-Muller codes for correcting a large number of errors will frequently correct more than $2^m - 1$ errors, and will always correct $2^m - 1$ or fewer errors.

These values are given for a 256 digit message in Table XV.

3. Each digit of the transmitted message is a parity check of a group of digits from the information source; the message cannot be broken down into information digits and check digits.

4. The decoding is accomplished by a number of majority decisions among different groups of message digits.

A technique will be described for using a Reed-Muller code efficiently to correct a number of small (± 1) errors for any code base b that is a multiple of 2, and also, at a small sacrifice of efficiency, a number of larger errors.

A theorem, stating that any code which is generated by a set of parity

checks will contain the same set of allowable messages as some systematic code, was proved by Hamming.¹¹ In particular, such a theorem indicates that a Reed-Muller code will contain the same set of allowable messages as some systematic code. This was also proved by Slepian,¹² who has given a simple method of deriving a systematic code generating the same set of messages as a Reed-Muller code. For convenience, such a code will be called an SERM code (Systematic Equivalent Reed-Muller code).

A Reed-Muller decoder serves to derive the information digits from a message in Reed-Muller code which may have been mutilated by noise. If a Reed-Muller decoder is followed by a Reed-Muller encoder, the combination serves as a noise eliminator (provided the noise is within the correction bounds of the code), since the output of the encoder is the noiseless Reed-Muller code message that is equivalent to the noisy message that entered the decoder. This property is useful since it means that any message, drawn from the set of Reed-Muller code messages, which has not been mutilated outside the bounds set up by a particular Reed-Muller code, will be restored to its original form, by a Reed-Muller decoder followed by a Reed-Muller encoder. Since an SERM code will produce only messages included in the set of messages of the corresponding Reed-Muller code, the SERM code can be used in conjunction with a Reed-Muller decoder and encoder to permit transmission over a noisy channel in a systematic code.

The two systems shown in Fig. 3 are therefore equivalent in their error correction properties. In both cases, messages from the set of Reed-Muller code messages are sent, and since the same decoder is used initially, both systems will correct errors in the received message in the same manner. The Reed-Muller encoder in the second system is required because a Reed-Muller decoder does not correct a message but derives information digits from the received message directly. The derived information digits, however, necessarily correspond to some corrected form of the received message and, in effect, the decoder performs the same correction as it would perform by deriving the corrected form of the message first.

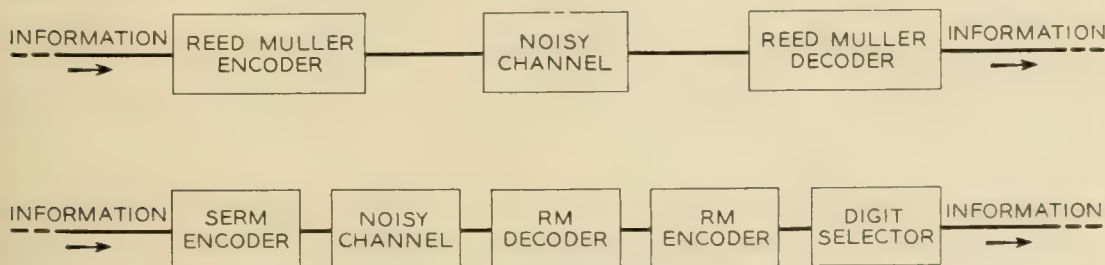


Fig. 3 — Equivalent systems using SERM and Reed-Muller codes.

TABLE XVI — MULTIPLE SMALL ERROR CORRECTION CODE
USING SERM CODES WITH DECIMAL CHANNEL

Message Length	Information Digits (Equivalent Decimal Digits)	Check Digits (Equivalent Decimal Digits)	No. of Small Errors Correctable per mes.
128	127.7	.3	0
128	125.3	2.7	1
128	116.9	11.1	3
128	100.0	28.0	7

This means that a Reed-Muller code can be adapted to the system shown in Fig. 2. The Systematic Binary Error Correction Code Encoder is simply an SERM encoder; this is permissible since the SERM codes are systematic. The Binary Decoder and Corrector is simply a Reed-Muller decoder followed by a Reed-Muller encoder. Everything else remains unchanged.

This scheme offers flexibility for the correction of large numbers of small errors. Proper initial error correction encoding of the original information digits will permit correction of a small number of large errors.

Table XVI shows some typical cases of the correction of many small errors in a decimal message as a function of the number of information and check digits in a message of constant length. For convenience, everything is shown in equivalent decimal digits, even though in the actual code, binary and quinary information digits are used. Only the first few entries are considered, since the message composed exclusively of the digits 0, 3, 6, 9 in which any number of small errors in a decimal channel may be corrected (this code is described by the first entry of Table V) is more efficient than the codes corresponding to subsequent entries on Table XVI. This code, which is very easy to instrument, will transmit the equivalent of 77 decimal digits in a 128 decimal digit message.

One problem not efficiently solved by these techniques is the multiple-error correction ternary channel problem. A technique which can be used is a code identical to the regular binary Reed-Muller Code, except that all equations will be modulo 3 instead of modulo 2. In decoding, this will sometimes require subtraction instead of addition; in modulo 2 equations there is no difference between these operations, but in modulo 3 equations, the two operations are distinct. The same procedure can be used for correcting multiple unrestricted errors in any base.

VII. ITERATIVE CODES

All the codes described above have one disadvantage; occasional excessive noise will yield a non-correctable message. In order to approach

error free transmission, some iterative coding procedure may be used. This problem has been solved by Elias.¹³ His methods are directly applicable to non-binary codes, since nothing restricts the digits to binary values.

In order to minimize the complexity of an iterative coding procedure, systematic codes are desirable. The advantages of the Reed-Muller code are significant however, especially for the case of a relatively noisy channel. A sound procedure for a binary channel would therefore be to use SERM codes, (see Fig. 3); such codes are more efficient than iterated Hamming Codes in a relatively noisy channel.

VIII. SUMMARY AND ANALYSIS

Many codes have been presented in this paper, all constructed by some combination of procedures involving linear congruence or modulo equations.

In most cases, more efficient codes exist. Exhaustive procedures exist for deriving maximum efficiency codes, although the codes derived in this manner usually require an extensive codebook, both at the encoder and at the decoder. Even for simple single error correction binary codes, the most efficient code is not always a systematic code. For example, the best systematic single error correction binary code working with an eight digit message has only 16 different allowable messages; it is known¹⁴ that a non-systematic code with at least 19 allowable messages exists.

In the case of non-binary codes, the situation is somewhat worse. Very few of the codes given in this paper take advantage of the fact that, for most situations, a digit that is incorrectly received as 0 or $b - 1$ is usually corrected only in one direction and no need exists to specify whether the correction is ± 1 . Most of the codes are arranged so that any received digit may be corrected either positively or negatively. No codes have been found which take full advantage of such a property, other than codebook codes, except for isolated instances of short message codes having symmetrical properties. For example, the single digit, single small error correction decimal code having 0, 3, 6, 9 as the allowable messages takes full advantage of this property, and is, at the same time, a true semi-systematic code.

It is extremely difficult to find the ultimate limits of efficiency of codebook codes. The exhaustive procedures are totally impractical except for very short messages. If an analysis is restricted to codes which do not take advantage of the property that certain values of digits may be corrected in only one direction, and it is assumed that each possible message is mutilated to the same number of incorrect messages, one

limit to the efficiency of codes may be found. This limit can be derived from the fact that an error correction code decoder and correction circuit must be able to convert any message which contains errors within the bounds of the correction performed by the code, into the value of the message as originally transmitted, or must be able to derive the original information which was fed into the encoder. Thus, if each message may be mutilated in w ways, and still be corrected, then at least w messages must be associated with each allowed message. This is indicated diagrammatically in Fig. 4. The messages produced by the encoder are shown at the left; each one fans out to $w - 1$ mutilated messages plus the original message. The decoder converts any of these w messages into the original message.

The value of w can be determined by taking all possible combinations of errors that can be corrected by a coding system. For example, for a code system which can correct up to $(d - 1)/2$ small errors in different digits in an n digit message, w is given by

$$w = \sum_{i=0}^{(d-1)/2} C_i^n 2^i, \quad (55)$$

where d is the minimum distance between messages, and

$$C_i^n = \frac{n!}{(n - i)! i!}.$$

This equation merely signifies that w is the sum of all combinations of positive and negative (accounting for the 2 term) errors in up to $(d - 1)/2$ different digits out of n digits. For single errors, $w = 2n + 1$.

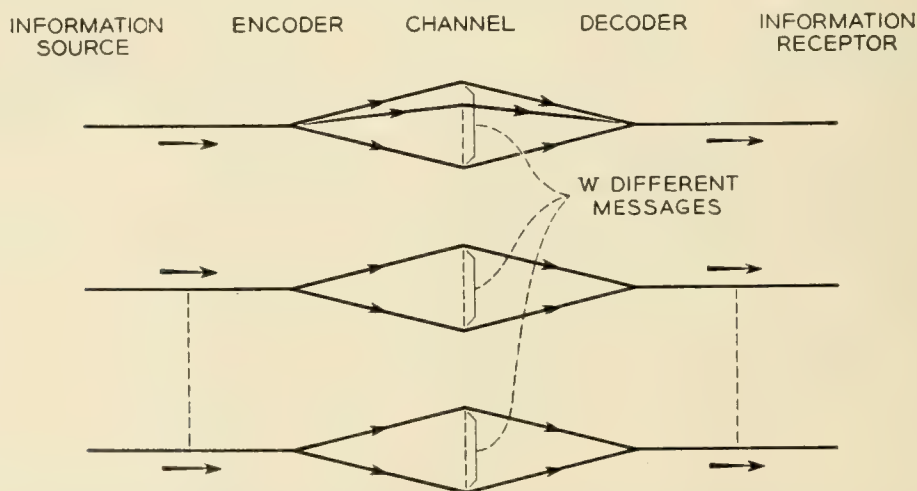


Fig. 4 — Graphical representation of an error correction code.

The number of different messages that can be produced by the encoder must be no greater than b^n/w , subject to the above restriction, b^n representing the maximum number of messages that the decoder may receive as an input. If only systematic and semi-systematic codes are considered, the number of messages is limited to multiples of powers of b and of the information component base β of mixed digits. The number of check states must be at least as large as w , so that w different correctors may be calculated and associated with w different corrections.

Subject to the above restrictions, the following statements may be made.

1. The systematic single small error correction codes derived using the rules of Section III are the most efficient systematic single small error correction codes possible. For those codes in which the two sides of inequality (8a) are equal, no code, not even a non-systematic code, is more efficient.

2. The systematic single unrestricted error correction codes derived using the rules of Section 4.1 are the most efficient systematic single unrestricted error correction codes. For those codes in which the two sides of inequality (31) are equal, no code is more efficient.

3. No codes are more efficient than those semi-systematic codes, derived using the rules of Section III, for which the two sides of inequalities (28) and (29) are equal and $m_1 = 0$. It is difficult to make more general statements about semi-systematic codes, because special techniques (such as those of Section VI), not all of which are known, may be used with these codes.

For multiple error correction codes, other techniques are both simpler and more efficient than the straight systematic and semi-systematic techniques described in Sections III, IV and V. One such scheme has been described in detail in Section VI. No codes have been found which approach the limit set by w , but the codes described in Section 6.2 are moderately efficient.

Throughout this paper, all techniques which involve vast complications at the expense of slight additional efficiency have been avoided. Codebook methods are always possible. If a technique is almost as complicated as a codebook technique with only slightly greater efficiency than a simple technique, the simple technique would always be used in practice, and the codebook satisfies the mathematical and theoretical requirements. In a sense, a really complicated technique is only useful for deriving a better lower limit for the maximum efficiency of a codebook code. In the non-binary case, however, a codebook system is considerably more efficient than any code system which does not take ad-

vantage of the fact that all transmitted messages are not mutilatable to an equal number of correctable received messages.

From the point of view of deriving lower limits to the maximum efficiency of a codebook technique, such a consideration is vital. Except for a few relatively trivial cases, no codes have been found which take significant advantage of the above consideration, for deriving such a limit.*

IX. CONCLUSION

In this paper, techniques have been presented for deriving error correction codes for non-binary systems. None of the methods presented are overly complicated, nor do they require excessive storage capacity for either the encoding or decoding and correction system.

The codes are sufficiently simple so that their use with a non-binary storage system may be considered, and the development of such a storage system should not be stopped because a system without flaws or not subject to noise cannot be realized.

An important disadvantage of using error correction codes with such a system is the time requirement. Correction usually requires a significant amount of time. This is probably one reason why the Hamming Code is not used more extensively. The more advanced and complicated codes, such as the Reed-Muller Codes, suffer particularly from the amount of time required for a correction. The codes described in this paper are therefore probably best suited to medium or low speed storages, which are not read too frequently.

A study of this type may be of some interest to those who have been considering the use of multi-state devices for building switching systems and computers, since this paper represents a study of a typical problem. Certain lessons may be derived from this study:

1. Restriction to a single number base for all operations is a severe handicap. The more advanced codes presented in this paper, require extensive use of different number base operations. The ability, inside the computer, to change number bases for different operations, may well be useful.

2. Different problems are best solved using different number bases. For example, the use of an even number base is desirable for multiple small error correction codes, while the use of a prime number base is desirable for correcting single large errors. It is the author's opinion that

* Note that this restriction has less significance in the case of binary codes. In a symmetrical channel with only two available signals, each value of a digit may be changed in as many ways, namely, one, as every other.

number bases which are the product of several small factors are best. Suggested values are six, ten and twelve. Number bases with two different prime factors, may offer an advantage, since they permit simple translation and change of number base among at least three different numbers.

In the comparison between binary and non-binary error correction codes, the following observations may be made:

1. Keeping the amount of information per message fixed, a binary single error correction code is less efficient than a non-binary single small error correction code, provided b , the channel base, is greater than three, but is more efficient than a non-binary single unrestricted error correction code.

2. Non-binary codes are slightly more complicated to implement than binary codes; this applies to multiple error correction codes as well as to single error correction codes. The amount of added complication is in no case really extensive.

It was initially hoped that this study might also produce some additional binary error correction techniques. One such technique was discovered: the use of a systematic equivalent Reed-Muller code to approach error free coding (see Section VII).

Finally, the author wishes to express the hope that further work on non-binary systems will be encouraged by this study.

ACKNOWLEDGEMENTS

This work was performed under the part-time Graduate Study Plan of Bell Telephone Laboratories at the Columbia University School of Engineering under the guidance of Prof. L. A. Zadeh. The author wishes to acknowledge the help of Prof. Zadeh, both in the selection of a dissertation topic and in the subsequent guidance of the study. In addition, a number of helpful discussions with C. Y. Lee and A. C. Rose helped to guide the research into a study of the most significant problems in the field.

REFERENCES

1. R. W. Hamming, Error Detecting and Error Correcting Codes, B.S.T.J., **29**, p. 147-160, April, 1950.
2. Of Current Interest, Elec. Engg., p. 871, Sept., 1956.
3. C. Y. Lee and W. H. Chen, Several-Valued Combinational Switching Circuits, Trans. A.I.E.E., **75**, Part I, p. 278-283, July, 1956.
4. D. Slepian, A Class of Binary Signaling Alphabets, B.S.T.J., **35**, p. 203-234, Jan., 1956.

5. Hamming, *op. cit.*, Section 1.
6. M. J. E. Golay, Notes on Digital Coding, Proc. I.R.E., **37**, p. 657, June, 1949.
7. W. Keister, A. E. Ritchie, and S. H. Washburn, *The Design of Switching Circuits*, D. Van Nostrand Co., Inc., New York, 1951, p. 316.
8. M. J. E. Golay, Binary Coding, 1954 Symposium on Information Theory, Trans. I.R.E., **PGIT-4**, p. 23-28, Sept., 1954.
9. G. H. Hardy and E. M. Wright, *An Introduction to the Theory of Numbers*, Oxford University Press, Oxford, 1954, p. 51, Theorem 57.
10. I. S. Reed, A Class of Multiple-Error-Correcting Codes and the Decoding Scheme, 1954 Symposium on Information Theory, Trans. I.R.E., **PGIT-4**, p. 38-49, Sept., 1954.
11. Hamming, *op. cit.*, Section 7.
12. D. Slepian, A Note on Two Binary Signaling Alphabets, Trans. I.R.E., **IT-2**, p. 84-86, June, 1956.
13. P. Elias, Error Free Coding, 1954 Symposium on Information Theory, Trans. I.R.E., **PGIT-4**, p. 29-38, Sept., 1954.
14. V. I. Siforov, On Noise Stability of a System with Error-Correcting Codes, Trans. I.R.E., **IT-2**, p. 109-115, Dec., 1956. See Table II, Column 8.

Shortest Connection Networks And Some Generalizations

By R. C. PRIM

(Manuscript received May 8, 1957)

The basic problem considered is that of interconnecting a given set of terminals with a shortest possible network of direct links. Simple and practical procedures are given for solving this problem both graphically and computationally. It develops that these procedures also provide solutions for a much broader class of problems, containing other examples of practical interest.

I. INTRODUCTION

A problem of inherent interest in the planning of large-scale communication, distribution and transportation networks also arises in connection with the current rate structure for Bell System leased-line services. It is the following:

Basic Problem — *Given a set of (point) terminals, connect them by a network of direct terminal-to-terminal links having the smallest possible total length (sum of the link lengths).* (A set of terminals is “connected,” of course, if and only if there is an unbroken chain of links between every two terminals in the set.) An example of such a *Shortest Connection Network* is shown in Fig. 1. The prescribed terminal set here consists of Washington and the forty-eight state capitals. The distances on the particular map used are accepted as true.

Two simple construction principles will be established below which provide simple, straight-forward and flexible procedures for solving the basic problem. Among the several alternative algorithms whose validity follows from the basic construction principles, one is particularly well adapted for automatic computation. The nature of the construction principles and of the demonstration of their validity leads quite naturally to the consideration, and solution, of a broad class of minimization problems comprising a non-trivial abstraction and generalization of the basic problem. This extended class of problems contains examples of practical

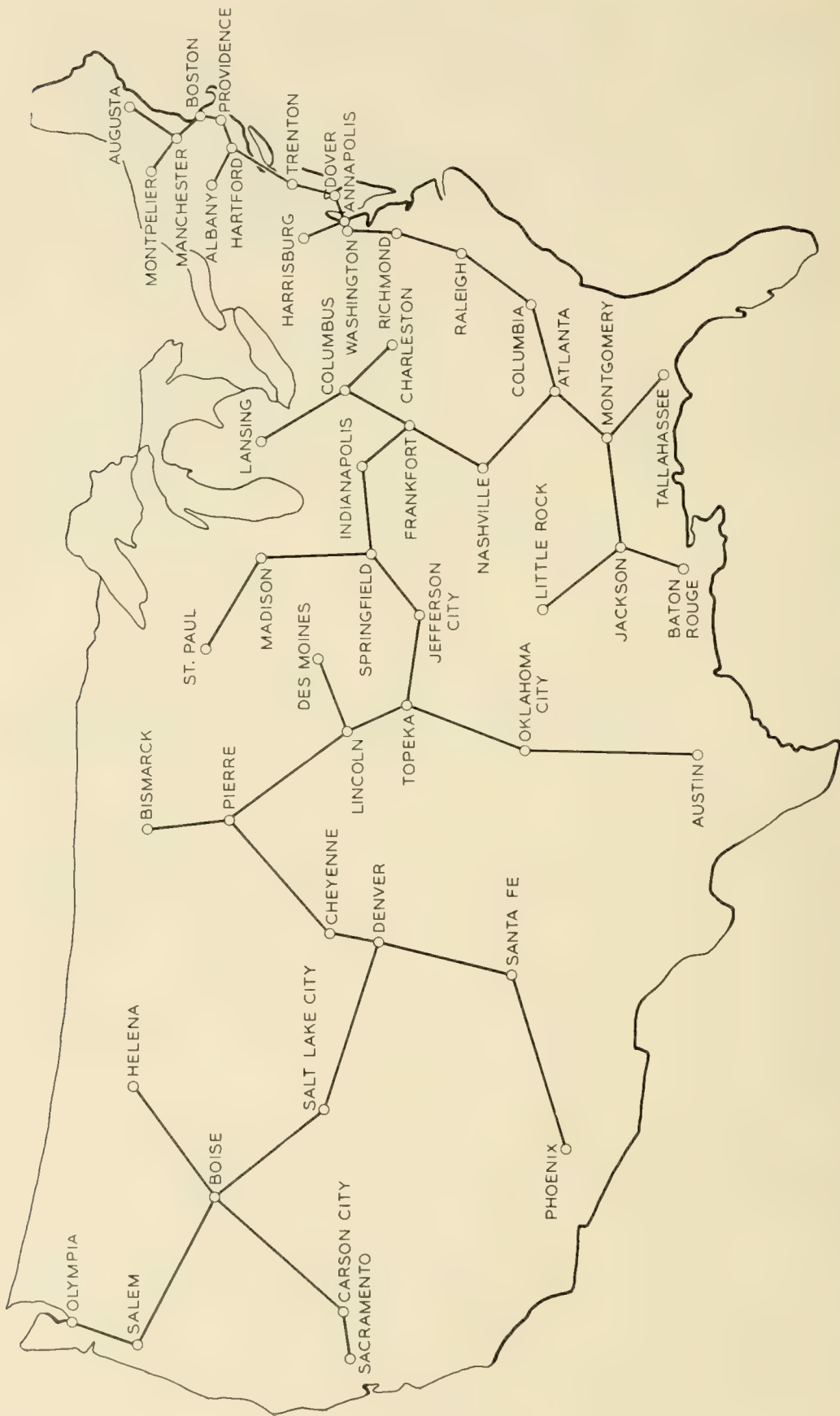


Fig. 1 — Example of a shortest connection network.

interest in quite different contexts from those in which the basic problem had its genesis.

II. CONSTRUCTION PRINCIPLES FOR SHORTEST CONNECTION NETWORKS

In order to state the rules for solution of the basic problem concisely, it is necessary to introduce a few, almost self-explanatory, terms. An *isolated terminal* is a terminal to which, at a given stage of the construction, no connections have yet been made. (In Fig. 2, terminals 2, 4, and 9 are the only isolated ones.) A *fragment* is a terminal subset connected by direct links, between members of the subset. (In Fig. 2, 8-3, 1-6-7-5, 5-6-7, and 1-6 are some of the fragments; 2-4, 4-8-3, 1-5-7, and 1-7 are

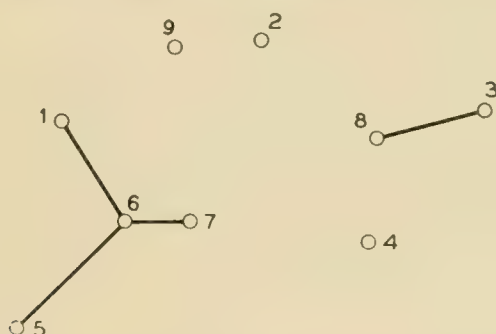


Fig. 2 — Partial connection network.

not fragments.) The *distance of a terminal from a fragment* of which it is not an element is the minimum of its distances from the individual terminals comprising the fragment. An *isolated fragment* is a fragment to which, at a given stage of the construction, no external connections have been made. (In Fig. 2, 8-3 and 1-6-7-5 are the only isolated fragments.) A *nearest neighbor of a terminal* is a terminal whose distance from the specified terminal is at least as small as that of any other. A *nearest neighbor of a fragment*, analogously, is a terminal whose distance from the specified fragment is at least as small as that of any other.

The two fundamental construction principles (P1 and P2) for shortest connection networks can now be stated as follows:

Principle 1 — *Any isolated terminal can be connected to a nearest neighbor.*

Principle 2 — *Any isolated fragment can be connected to a nearest neighbor by a shortest available link.*

For example, the next steps in the incomplete construction of Fig. 2 could be any one of the following:

- (1) add link 9-2 (P1 applied to Term. 9)

- (2) add link 2-9 (P1 applied to Term. 2)
- (3) add link 4-8 (P1 applied to Term. 4)
- (4) add link 8-4 (P2 applied to frag. 3-8)
- (5) add link 1-9 (P2 applied to frag. 1-6-7-5).

One possible sequence for completing this construction is: 4-8 (P1), 8-2 (P2), 9-2 (P1), and 1-9 (P2). Another is: 1-9 (P2), 9-2 (P2), 2-8 (P2), and 8-4 (P2).

As a second example, the construction of the network of Fig. 1 could have proceeded as follows: Olympia-Salem (P1), Salem-Boise (P2), Boise-Salt Lake City (P2), Helena-Boise (P1), Sacramento-Carson City (P1), Carson City-Boise (P2), Salt Lake City-Denver (P2), Phoenix-Santa Fe (P1), Santa Fe-Denver (P2), and so on.

The kind of intermixture of applications of P1 and P2 demonstrated here is very efficient when the shortest connection network is actually being laid out on a map on which the given terminal set is plotted to scale. With only a few minutes of practice, an example as complex as that of Fig. 1 can be solved in less than 10 minutes. Another mode of procedure, making less use of the flexibility permitted by the construction principles, involves using P1 only once to produce a single fragment, which is then extended by successive applications of P2 until the network is completed. This highly systematic variant, as will emerge later, has advantages for computer mechanization of the solution process. As applied to the example of Fig. 1, this algorithm would proceed as follows if Sacramento were the indicated initial terminal: Sacramento-Carson City, Carson City-Boise, Boise-Salt Lake City, Boise-Helena, Boise-Salem, Salem-Olympia, Salt Lake City-Denver, Denver-Cheyenne, Denver-Santa Fe, and so on.

Since each application of either P1 or P2 reduces the total number of isolated terminals and fragments by one, it is evident that an N -terminal network is connected by $N-1$ applications.

III. VALIDATION OF CONSTRUCTION PRINCIPLES

The validity of P1 and P2 depends essentially on the establishment of two *necessary* conditions (NC1 and NC2) for a *shortest* connection network (SCN):

Necessary Condition 1 — *Every terminal in a SCN is directly connected to at least one nearest neighbor.*

Necessary Condition 2 — *Every fragment in a SCN is connected to at least one nearest neighbor by a shortest available path.*

To simplify the argument, it will at first be assumed that all distances

between terminals are different, so that each terminal or fragment has a single, uniquely defined, nearest neighbor. This restriction will be removed later.

Consider first NC1. Suppose there is a SCN for which it is untrue. Then [Fig. 3(a)] some terminal, t , in this network is not directly joined to its nearest neighbor, n . Since the network is connected, t is necessarily joined directly to some one or more terminals other than n — say $f_1, \dots, f_r$. For the same reason, n is necessarily joined through some chain, C , of one or more links to one of $f_1, \dots, f_r$ — say to f_k . Now if the link $t - f_k$ is removed from the network and the link $t - n$ is added [Fig. 3(b)], the connectedness of the network is clearly not destroyed — f_k being joined to t through n and C , rather than directly. However, the total length of the network has now been *decreased*, because, by hypothesis, $t - n$ is shorter than $t - f_k$. Hence, contrary to the initial supposition, the network contradicting NC1 could not have been the *shortest*, and the truth of NC1 follows. From NC1 follows P1, which merely *permits* the addition of links which NC1 shows *have* to appear in the final SCN.

Turning now to NC2, the above argument carries over directly if t is thought of as a fragment of the supposed contradictory SCN, rather than as an individual terminal — provided, of course that the $t - n$ link substituted for $t - f_k$ is the shortest link from n to any of the terminals belonging to t . From the validity of NC2 follows P2 — again the links whose addition is permitted by P2 are all necessary, by NC2, in the final SCN.

The temporary restrictive assumption that no two inter-terminal distances are identical must now be removed. A reappraisal of the proofs of NC1 and NC2 shows that they are still valid if n is not the only terminal at distance $t - n$ from t , for in the supposedly contradictory network the distance $t - f_k$ must be *greater* than $t - n$. What remains to be established is that the length of the final connection network resulting

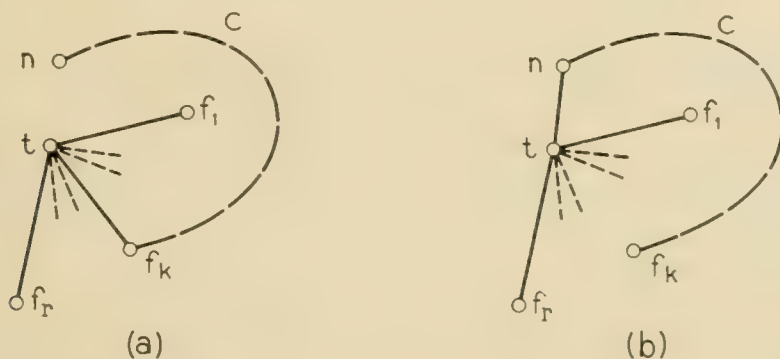


Fig. 3 — Schematic demonstration of NC1.

from successive applications ($N - 1$ for N terminals) of P1 and P2 is independent of *which* nearest neighbor is chosen for connection at a stage when more than one nearest neighbor to an isolated terminal or t is available. This is a consequence of the following considerations: For a prescribed terminal set there are only a *finite* number of connection networks (certainly fewer than $C_{N-1}^{N(N-1)/2}$ — the number of distinct ways of choosing $N - 1$ links from the total of $N(N - 1)/2$ possible links). The length of each one of this finite set of connection networks is a *continuous* function of the individual interterminal distances, d_{ij} (as a matter of fact, it is a linear function with coefficients 0 and 1). With the d_{ij} specified, the length, L , of a shortest connection network is simply the smallest length in this finite set of connection network lengths. Therefore L is uniquely determined. (It may, of course, be associated with more than one of the connection networks.) Now, if at each stage of construction employing P1 and P2 at which a choice is to be made among two or more nearest neighbors $n_1, \dots, n_r$ of an isolated terminal (or fragment) t , a small positive quantity, ϵ , is subtracted from any specific one of the distances $d_{tn_1}, \dots, d_{tn_r}$ — say from d_{tn_k} — the construction will be uniquely determined. The total length, L' , of the resulting SCN for the modified problem will now depend on ϵ , as well as on the d_{ij} of the original terminal set. The dependence on ϵ will be *continuous*, however, because the minimum of a finite set of continuous functions of ϵ (the set of lengths of all connection networks of the modified problem) is itself a continuous function of ϵ . Hence, as ϵ is made vanishingly small, L' approaches L , regardless of which “nearest neighbor” links were chosen for shortening to decide the construction.

IV. ABSTRACTION AND GENERALIZATION

In the examples of Figs. 1 and 2, the terminal set to be connected was represented by points on a distance-true map. The decisions involved in applying P1 and P2 could then be based on visual judgements of relative distances, perhaps augmented by application of a pair of dividers in a few close instances. These direct geometric comparisons can of course, be replaced by numerical ones if the inter-terminal distances are entered on the terminal plot, as in Fig. 4(a). The application of P1 and P2 goes through as before, with the relevant “nearest neighbors” determined by a comparison of numerical labels, rather than by a geometric scanning process. For example, P1 applied to Terminal 5 of Fig. 4(a) yields the Link 5-6 of the SCN of Fig. 4(b), because 4.6 is less than 5.6, 8.0, 8.5, and 5.1, and so on.

When the construction of shortest connection networks is thus reduced to processes involving only the numerical distance labels on the various possible links, the *actual location* of the points representing the various terminals in a graphical representation of the problem is, of course, inconsequential. The problem of Fig. 4(a) can just as well be represented by Fig. 5(a), for example, and P1 and P2 applied to obtain the SCN of Fig. 5(b). The original metric problem concerning a set of points in the plane has now been abstracted into a problem concerning *labelled graphs*. The correspondence between the terminology employed thus far and more conventional language of Graph Theory is as follows:

terminal $\leftrightarrow$ vertex

possible link $\leftrightarrow$ edge

length of link $\leftrightarrow$ "length" (or "weight") of edge

connection network $\leftrightarrow$ spanning subgraph

(without closed loops) $\leftrightarrow$ (spanning subtree)

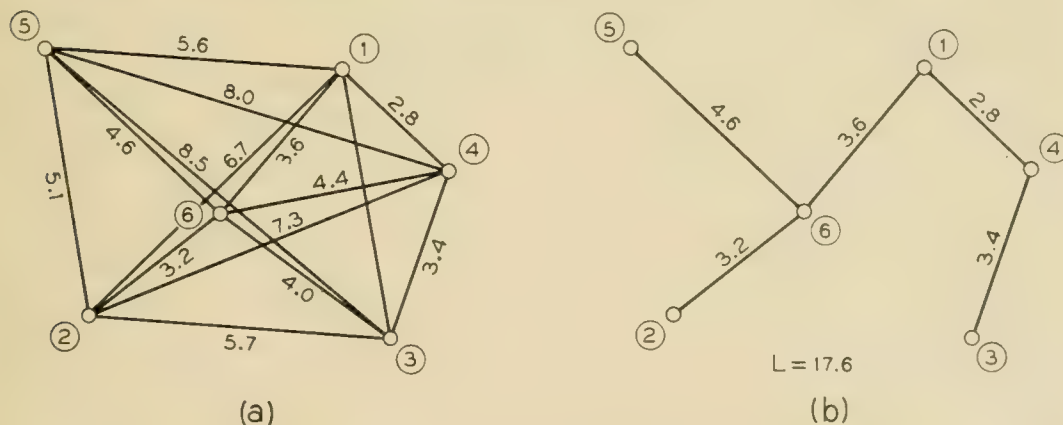


Fig. 4 — Example of a shortest spanning subtree of a complete labelled graph.

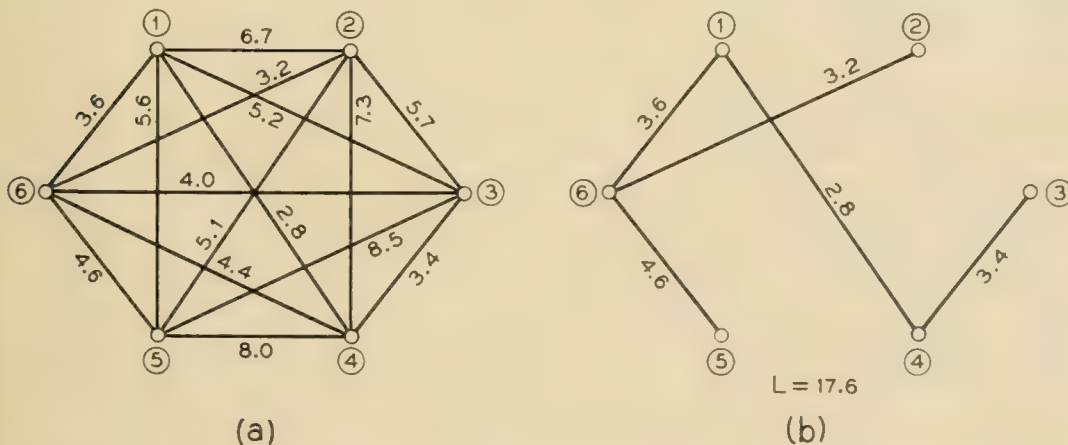


Fig. 5 — Example of a shortest spanning subtree of a complete labelled graph.

shortest connection network $\leftrightarrow$ shortest spanning subtree
 SCN $\leftrightarrow$ SSS

It will be useful and worthwhile to carry over the concepts of "fragment" and "nearest neighbor" into the graph theoretic framework. P1 and P2 can now be regarded as construction principles for finding a shortest spanning subtree of a labelled graph.

In the originating context of connection networks, the graphs from which a shortest spanning subtree is to be extracted are *complete graphs*; that is, graphs having an edge between every pair of vertices. It is natural, now, to generalize the original problem by seeking shortest spanning subtrees for arbitrary connected labelled graphs. Consider, for example, the labelled graph of Fig. 6(a) which is derived from that of Fig. 5(a) by deleting some of the edges. (In the connection network context, this is equivalent to barring direct connections between certain terminal pairs.) It is easily verified that NC1 and NC2, and hence P1 and P2, are valid also in these more general cases. For the example of Fig. 6(a), they yield readily the SSS of Fig. 6(b).

As a further generalization, it is not at all necessary for the validity of P1 and P2 that the edge "lengths" in the given labelled graph be derived, as were those of Figs. 4–6, from the inter-point distances of some point set in the plane. *P1 and P2 will provide a SSS for any connected labelled graph with any set of real edge "lengths."* The "lengths" need not even be positive, or of the same sign. See, for example, the labelled graph of Fig. 7(a) and its SSS, Fig. 7(b). It might be noted in passing that this degree of generality is sufficient to include, among other things, shortest connection networks in an arbitrary number of dimensions.

A further extension of the range of problems solved by P1 and P2 follows trivially from the observation that the *maximum* of a set of

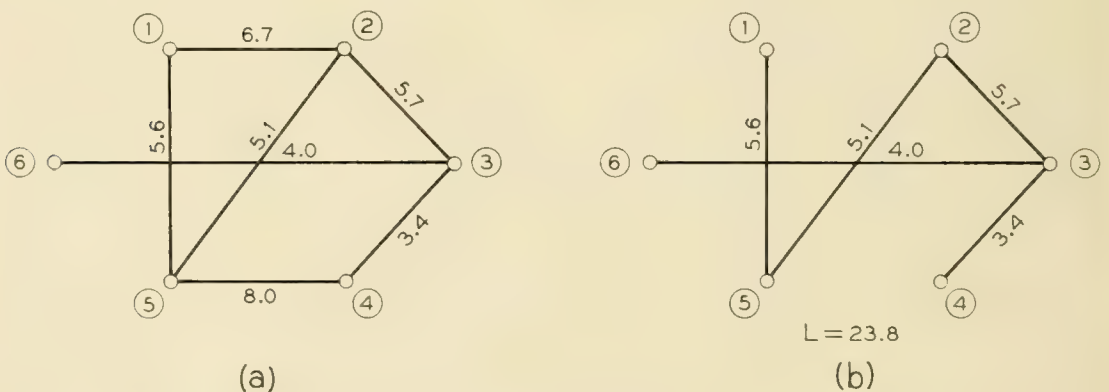


Fig. 6 — Example of a shortest spanning subtree of an incomplete labelled graph.

real numbers is the same as the negative of the *minimum* of the negatives of the set. Therefore, P1 and P2 can be used to construct a *longest spanning subtree* by changing the signs of the “lengths” on the given labelled graph. Fig. 8 gives, as an example, the longest spanning subtree for the labelled graph of Figs. 4(a) and 5(a).

It is easy to extend the arguments in support of NC1 and NC2 from the simple case of minimizing the sum to the more general problems of minimizing an arbitrary increasing symmetric function, or of maximizing an arbitrary decreasing symmetric function, of the edge “lengths” of a spanning subtree. (A symmetric function of n variables is one whose value is unchanged by any interchanges of the variable values; e.g., $x_1 + x_2 + \cdots + x_n$, $x_1 x_2 \cdots x_n$, $\sin 2x_1 + \sin 2x_2 + \cdots + \sin 2x_n$, $(x_1^3 + x_2^3 + \cdots + x_n^3)^{1/2}$, etc.) From this follow the non-trivial generalizations.

The *shortest spanning subtree* of a connected labelled graph also minimizes all increasing symmetric functions, and maximizes all decreasing symmetric functions, of the edge “lengths.”

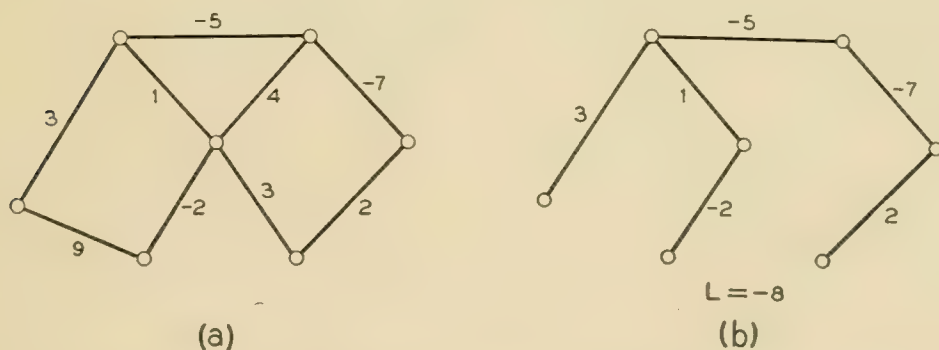


Fig. 7 — Example of a shortest spanning subtree of a labelled graph with some “lengths” negative.

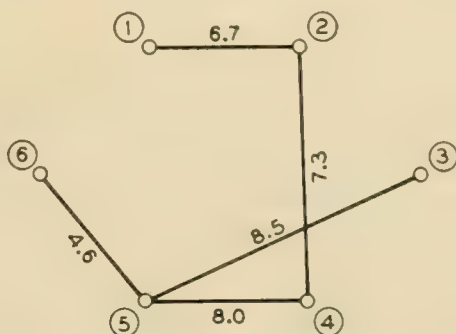


Fig. 8 — The longest spanning subtree of the labeled graph of Figs. 4(a) and 5(a).

The *longest spanning subtree* of a connected labelled graph also maximizes all increasing symmetric functions, and minimizes all decreasing symmetric functions, of the edge "lengths."

For example, with positive "lengths" the same spanning subtree that minimizes the *sum* of the edge "lengths" also minimizes the *product* and the *square root of the sum of the squares*. At the same time, it maximizes the sum of the reciprocals and the product of the arc cotangents.

It seems likely that these extensions of the original class of problems soluble by P1 and P2 contain many examples of practical interest in quite different contexts from the original connection networks. A not entirely facetious example is the following: A message is to be passed to all members of a certain underground organization. Each member knows some of the other members and has procedures for arranging a rendezvous with anyone he knows. Associated with each such possible rendezvous — say between member "*i*" and member "*j*" — is a certain probability, p_{ij} , that the message will fall into hostile hands. How is the message to be distributed so as to minimize the over-all chances of its being compromised? If members are represented as vertices, possible rendezvous as edges, and compromise probabilities as "length" labels in a labelled graph, the problem is to find a spanning subtree for which $1 - \prod(1 - p_{ij})$ is minimized. Since this is an increasing symmetric function of the p_{ij} 's, this is the same as the spanning subtree minimizing $\sum p_{ij}$, and this is easily found by P1 and P2.

Another application, closer to the original one, is that of minimizing the lengths of wire used in cabling panels of electrical equipment. Restrictions on the permitted wiring patterns lead to shortest connection network problems in which the effective distances between terminals are not the direct terminal-to-terminal distances. Thus the more general viewpoint of the present section is applicable.

V. COMPUTATIONAL TECHNIQUE

Return now to the exemplary shortest connection network problem of Figs. 4(a) and 5(a) which served as the center for discussion of the arithmetizing of the metric factors in the Basic Problem. It is evident that the actual drawing and labelling of all the edges of a complete graph will get very cumbersome as the number of vertices increases — an N -vertex graph has $(1/2)(N^2 - N)$ edges. For large N , it is convenient to organize the numerical metric information in the form of a *distance table*, such as Fig. 9, which is equivalent in content to Fig. 4(a) or Fig.

5(a). (A distance table can also be prepared to represent an incomplete labelled graph by entering the length of non-existent edges as ∞ .)

When it is desired to determine a shortest connection network directly from the distance table representation — either manually, or by machine computation — one of the numerous particular algorithms obtainable

	1	2	3	4	5	6
1	—	6.7	5.2	2.8	5.6	3.6
2	6.7	—	5.7	7.3	5.1	3.2
3	5.2	5.7	—	3.4	8.5	4.0
4	2.8	7.3	3.4	—	8.0	4.4
5	5.6	5.1	8.5	8.0	—	4.6
6	3.6	3.2	4.0	4.4	4.6	—

Fig. 9 — Distance table equivalent of Figs. 4(a) and 5(a).

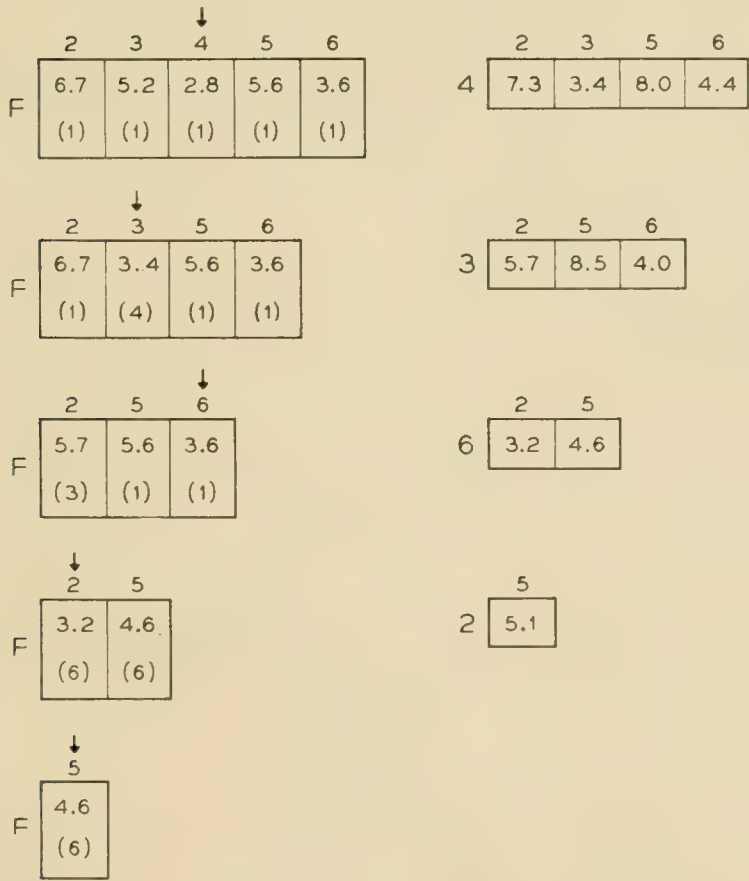


Fig. 10 — Illustrative computational determination of a shortest connection network.

by restricting the freedom of choice allowed by P1 and P2 is distinctly superior to other alternatives. This variant is the one in which P1 is used but once to produce a single isolated fragment, which is then extended by repeated applications of P2.

The successive steps of an efficient computational procedure, as applied to the example of Fig. 9, are shown in Fig. 10. The entries in the top rows of the successive F tables are the distances from *the* connected fragment to the unconnected terminals at each stage of fragment growth. The entries in parentheses in the second rows of these tables indicate the nearest neighbor *in the fragment* of the external terminal in question. The computation is started by entering the first row of the distance table into the F table (to start the growing fragment from Terminal 1). A smallest entry in the F table is then selected (in this case, 2.8, associated with External Terminal 4 and Internal Terminal 1). The link 1-4 is deleted from the F table and entered in the Solution Summary (Fig. 11). The remaining entries in the first stage F table are then compared with the corresponding entries in the “4” row of the distance table (reproduced beside the first F table). If any entry of this “added terminal” distance table is smaller than the corresponding F table entry, it is substituted for it, with a corresponding change in the parenthesized index. (Since 3.4 is less than 5.2, the 3 column of the F table becomes 3.4/(4).) This process is repeated to yield the list of successive nearest neighbors to the growing fragment, as entered in Fig. 11. The F and “added terminal” distance tables grow shorter as the number of unconnected terminals is decreased.

This computational procedure is easily programmed for an automatic computer so as to handle quite large-scale problems. One of its advantages is its avoidance of checks for closed cycles and connectedness. Another is that it never requires access to more than two rows of distance data at a time — no matter how large the problem.

SOLUTION SUMMARY	
LINK	LENGTH
1 - 4	2.8
4 - 3	3.4
1 - 6	3.6
6 - 2	3.2
6 - 5	4.6

Fig. 11 — Solution summary for computation of Fig. 10.

VI. RELATED LITERATURE AND PROBLEMS

J. B. Kruskal, Jr.¹ discusses the problem of constructing shortest spanning subtrees for labelled graphs. He considers only distinct and positive sets of edge lengths, and is primarily interested in establishing uniqueness under these conditions. (This follows immediately from NC1 and NC2.) He also, however, gives three different constructions, or algorithms, for finding SSS's. Two of these are contained as special cases in P1 — P2. The third is — “Perform the following step as many times as possible: Among the edges not yet chosen, choose the longest edge whose removal will not disconnect them” While this is not directly a special case of P1 — P2, it is an obvious corollary of the special case in which the *shortest* of the edges permitted by P1 — P2 is selected at each stage. Kruskal refers to an obscure Czech paper² as giving a construction and uniqueness proof inferior to his.

The simplicity and power of the solution afforded by P1 and P2 for the Basic Problem of the present paper comes as something of a surprise, because there are well-known problems which *seem* quite similar in nature for which no efficient solution procedure is known.

One of these is *Steiner's Problem*: Find a shortest connection network for a given terminal set, with freedom to add additional terminals wherever desired. A number of necessary properties of these networks are known³ but do not lead to an effective solution procedure.

Another is the *Traveling Salesman Problem*: Find a closed path of minimum length connecting a prescribed terminal set. Nothing even approaching an effective solution procedure for this problem is now known (early 1957).

REFERENCES

1. J. B. Kruskal, Jr., On the Shortest Spanning Subtree of a Graph and the Traveling Salesman Problem, Proc. Amer. Math. Soc. **7**, pp. 48-50, 1956.
2. Otakar Borůvka, On a Minimal Problem, Práce Moravské Pridovedecké Společnosti, **3**, 1926.
3. R. Courant and H. Robbins, *What is Mathematics*, 4th edition, Oxford Univ. Press, N. Y., 1941, pp. 374 *et seq.*

A Network Containing a Periodically Operated Switch Solved by Successive Approximations

By C. A. DESOER

(Manuscript received June 15, 1956)

This paper concerns itself with the analysis of a type of periodically switched network that might be used in time multiplex systems. The economics of the situation require that the ratio of the switch closure time τ to the switching period T be small. Using this assumption, the analysis is performed by successive approximations. More precisely the zeroth approximation to the transmission is obtained from a block diagram analogous to those used in sampled servomechanisms. From the convergence proof of the successive approximation scheme, it follows that when τ/T is small, the zeroth approximation is very close to the exact transmission. A discussion of some examples is included.

CONTENTS		Page
I.	Introduction	1404
II.	Description of the system	1404
III.	Method of solution	1407
IV.	The zeroth approximation	1408
	4.1 Introduction	1408
	4.2 Description of the block diagram	1410
	4.3 Analysis of the block diagram	1411
V.	Transmission loss	1413
VI.	A simple example	1414
VII.	Numerical examples	1416
VIII.	The successive approximation scheme	1417
	8.1 Preliminary steps	1418
	8.2 Matrix description of the successive approximations	1420
	8.3 Convergence proof	1421
IX.	Modification of the block diagram to improve the zeroth approximation	1422
X.	Conclusions	1423
Appendices		
	I. Analysis of the resonant circuit	1423
	II. Study of the limiting case $T \rightarrow 0$	1425
	III. Zeroth approximation in the case where N_1 is not identical to N_2	1425
	IV. Derivation of equation (24)	1426

I. INTRODUCTION

One main contributor to the cost of transmission circuits is the transmission medium itself. Thus it is important to share the transmission medium among as many messages as possible. One possible method is the frequency multiplex where each message utilizes a different frequency band of the whole band available in the medium. An alternate method is the time multiplex where each message is assigned a time slot of duration τ and has access to that time slot once every T seconds. It is obvious that the economics of the situation requires that τ be as small as possible and T as large as possible so that the largest possible number of messages are transmitted over the medium. For this very reason the analysis of periodically switched networks is of special interest in the case where τ/T is small.

W. R. Bennett⁴ has published an exact analysis of this problem without any restrictions either on the network or on the ratio τ/T . It is believed, however, that the analysis presented in this paper will, in most practical cases, give the desired answer with a considerable reduction in the amount of calculations. The simplification of the analysis is mainly a result of the assumption that τ/T is small.

First the successive approximation method of solution will be discussed in general terms. Next it will be shown that the zeroth approximation to the transmission through the network can be obtained from the gain of a block diagram analogous to those used in the analysis of sampled servomechanisms. The nature of the zeroth approximation is further clarified by some general discussion and some examples. Next it is shown that the successive approximations converge. The convergence proof then suggests some slight modifications of the block diagram to obtain a more accurate solution.

II. DESCRIPTION OF THE SYSTEM

The system under consideration is shown on Fig. 1. It consists of two reactive networks N_1 and N_2 connected through a switch S which is itself in series with an inductance ℓ . N_1 is driven at its terminal pair (1)

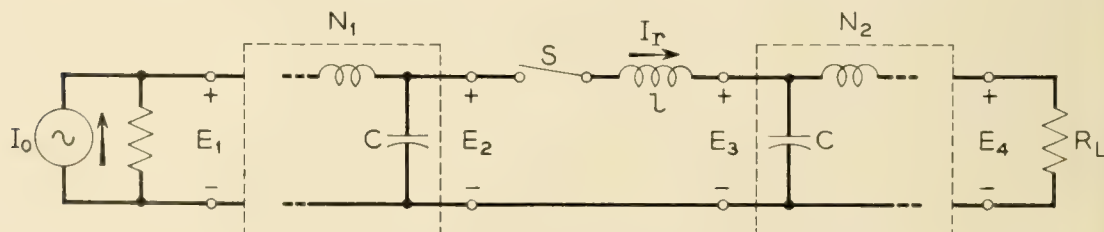
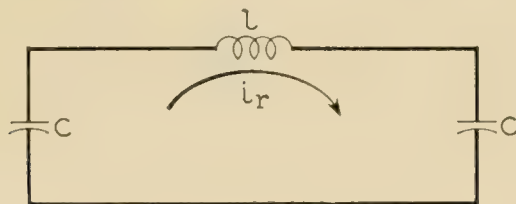


Fig. 1 — System under consideration.



$$f_r = \frac{1}{2\pi\sqrt{L\frac{C}{2}}} = \frac{\omega_0}{2\pi} \quad \tau = \frac{1}{2f_r} = \frac{\pi}{\omega_0} = \pi\sqrt{L\frac{C}{2}}$$

Fig. 2 — Resonant circuit.

by a current source I_0 which is shunted by a one ohm resistor. N_2 is also terminated at its terminal pair (1) by a one ohm resistor R_L which is the load resistor of the system. The switch S is periodically closed for a duration τ . The switching period is T . Thus if the switch is closed during the interval $(0, \tau)$ it will be closed during the intervals $(nT, nT + \tau)$ for $n = 1, 2, 3, \dots$. The inductance ℓ is selected so that the series circuit shown on Fig. 2 has a resonant frequency $f_r = 1/2\tau$; i.e., the time τ during which the switch is closed is exactly one-half period of the circuit of Fig. 2.

The switch S acts as a sampler and, as a result of the well-known modulating properties of sampled systems, the sampling period T must be chosen such that the frequency $1/2T$ is larger than any of the frequencies present in the signals generated by I_0 . Furthermore, in order to eliminate all the sidebands generated by the switching, N_2 must have a high insertion loss for all frequencies above $1/2T$ cps.

In the analysis that follows networks N_1 and N_2 will be assumed to be identical: it should, however, be stressed that this assumption is not necessary for the proposed method of analysis.* This assumption is made because in the practical situation which motivated this analysis N_1 and N_2 were identical since transmission in both directions was required.

In order for the system under consideration to achieve the maximum degree of multiplexing, the closure time τ of the switch will be taken as small as practically possible and the switching period T as large as possible (consistent with the bandwidth of the signals to be transmitted). As a result the ratio τ/T is very small, of the order of 10^{-2} or less in practical cases. Consequently the resonant frequency f_r of the series resonant circuit shown on Fig. 2, is many times larger than any of the natural frequencies of N_1 and N_2 .

* The modifications required for the case where N_1 is not identical to N_2 are given in Appendix IV.

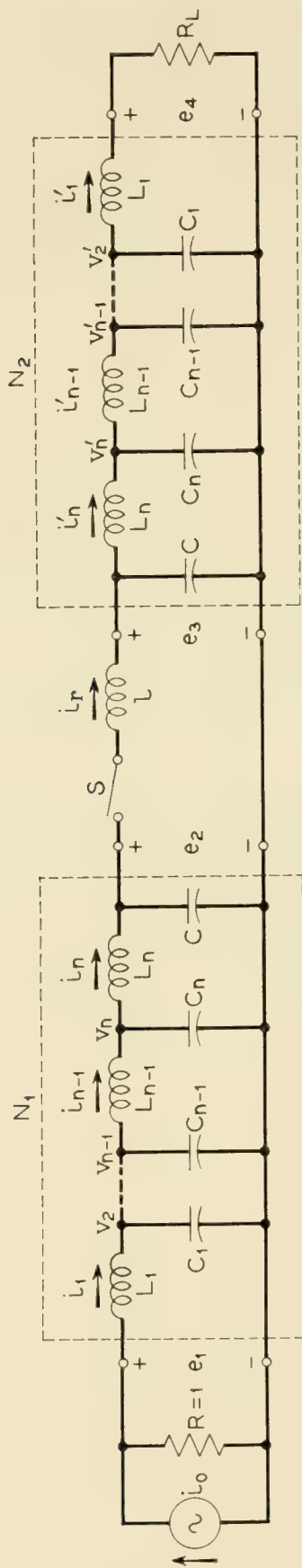


Fig. 3 — System under consideration when N_1 and N_2 are lossless ladders.

The problem is to determine the relation between E_4 , the voltage across R_L , and I_0 .

III. METHOD OF SOLUTION

Let us first write the equations of the system. Obviously the equations will depend on the exact configuration of the networks N_1 and N_2 . For simplicity we shall write them for the case where N_1 and N_2 are dissipationless low-pass ladder networks. As will become apparent later this assumption is not essential to the argument. What is essential, however, is the fact that both N_1 and N_2 should start (looking in from the switch) with a shunt capacitor C and a series inductance L_n , the element value of L_n being much larger than ℓ . Using a method of analysis advocated by T. R. Bashkow,⁵ we obtain, for the network of Fig. 3, the equations:

$$\left. \begin{aligned} L_1 \frac{di_1}{dt} &= -Ri_1 - v_2 + Ri_0 \\ C_1 \frac{dv_2}{dt} &= i_1 - i_2 \\ &\vdots \\ C_n \frac{dv_n}{dt} &= i_{n-1} - i_n \\ L_n \frac{di_n}{dt} &= v_n - e_2 \end{aligned} \right\} I_a \quad (1.a)$$

$$\left. \begin{aligned} C \frac{de_2}{dt} &= i_n - i_r \Delta(t) \quad (1.b) \\ \ell \frac{di_r}{dt} &= [e_2 - e_3] \Delta(t) \quad (1.c) \\ C \frac{de_3}{dt} &= i_r \Delta(t) - i_n' \quad (1.d) \end{aligned} \right\} R \quad (1)$$

$$\left. \begin{aligned} L_n \frac{di_n'}{dt} &= e_3 - v_n' \\ C_n \frac{dv_n'}{dt} &= i_n' - i_{n-1}' \\ &\vdots \\ C_1 \frac{dv_2'}{dt} &= i_2' - i_1' \\ L_1 \frac{di_1'}{dt} &= v_2' - R_L i_1' \end{aligned} \right\} I_b$$

where

$$\Delta(t) = \sum_{k=-\infty}^{+\infty} [u(t - kT) - u(t - kT - \tau)], \quad (2)$$

with $u(t) = 1$ for $t > 0$, and $u(t) = 0$ for $t < 0$.

This system of linear time varying equations may be broken up into three sub-systems I_a , R and I_b . It is this subdivision that suggests a successive approximation scheme that will be shown to converge to the exact solution.

The zeroth approximation is obtained as follows: when the switch is closed, i.e., $\Delta(t) = 1$, the resonant current i_r is much larger than the currents i_n and i_n' . Thus, during the switch closure time, i_n and i_n' are neglected with respect to i_r in (1.b) and (1.d). Hence when $\Delta(t) = 1$ the system R may be solved for $i_r(t)$, $e_2(t)$ and $e_3(t)$ in terms of the initial conditions. The resulting function $e_2(t)$ and given function $i_0(t)$ are then the forcing functions of the system I_a . The other function $e_3(t)$ is the forcing function of the system I_b . Under these assumptions, the periodic steady-state solution corresponding to an applied current $i_0(t) = I_0 e^{i\omega t}$ is easily obtained.

The zeroth approximation will be distinguished by a subscript "0". Thus $i_{r0}(t)$ is the (steady state) zeroth approximation to the exact solution $i_r(t)$.

The first approximation will be the solution of the system (1), provided that during the switch closure time the functions $i_n(t)$ and $i_n'(t)$ in (1.b) and (1.d) are respectively replaced by the known functions $i_{n0}(t)$ and $i_{n0}'(t)$. And, more generally, the $(k + 1)$ th approximation will be the solution of (1) provided that during the switch closure time, the functions $i_n(t)$ and $i_n'(t)$, in (1.b) and (1.d), are respectively replaced by the known solutions for $i_n(t)$, and $i_n'(t)$ given by the k th approximation. It will be shown later that this successive approximation scheme converges. Let us first describe a simple method for obtaining the zeroth approximation.

IV. THE ZEROTH APPROXIMATION

4.1 Introduction

The problem is to obtain the steady-state solution of (1) under the excitation $i_0(t) = I_0 e^{i\omega t}$. Using the approximations indicated above, during the switch closure time (that is when $\Delta(t) = 1$) the system R becomes

$$C \frac{de_2}{dt} = -i_r(t)\Delta(t), \quad (3)$$

$$\ell \frac{di_r}{dt} = [e_2 - e_3]\Delta(t), \quad (4)$$

$$C \frac{de_3}{dt} = i_r(t)\Delta(t). \quad (5)$$

Differentiating the middle equation and eliminating de_2/dt and de_3/dt we get for $0 \leq t < \tau$:

$$\frac{d^2 i_r}{dt^2} = -\frac{2}{\ell C} i_r \Delta(t) + \frac{1}{\ell} [e_2(t) - e_3(t)]\delta(t) \quad (6)$$

in which we used the notation $\delta(t)$ for the Dirac function and the knowledge that

$$\frac{d\Delta(t)}{dt} = \delta(t) - \delta(t - \tau). \quad (7)$$

Equation (6) represents the behavior of the resonant circuit of Fig. 2 for the following initial conditions:

$$i_r(0+) = 0, \quad (8)$$

$$\frac{di_r(0+)}{dt} = \frac{e_2(0) - e_3(0)}{\ell}. \quad (9)$$

In Appendix I it is shown that the resulting current $i_r(t)$ is, for the interval $0 \leq t < \tau$,

$$i_r(t) = C[e_2(0) - e_3(0)]s_1(t), \quad (10)$$

where

$$s_1(t) = \begin{cases} \frac{\pi}{2\tau} \sin \frac{\pi t}{\tau} = \frac{1}{2} \omega_0 \sin \omega_0 t & \text{for } 0 \leq t < \tau \\ 0 & \text{elsewhere} \end{cases} \quad (11)$$

with

$$\omega_0 = \frac{\pi}{\tau} = \sqrt{\frac{2}{\ell C}}. \quad (12)$$

Thus the zeroth approximation to the exact $i_r(t)$ is given for the interval $0 \leq t \leq T$ by

$$i_{r0}(t) = C[e_2(0) - e_3(0)]s_1(t). \quad (13)$$

We shall now show that the zeroth approximation may be conveniently obtained from the block diagram of Fig. 4.

4.2 Description of the Block Diagram

All the blocks of the block diagram are unilateral and their corresponding transfer functions are defined in the following. Capital symbols represent $\mathcal{L}$ -transform of the corresponding time functions, thus $I_0(p)$ is the $\mathcal{L}$ -transform of $i_0(t)$.

Referring to Fig. 1,

$$Z_{12}(p) = \left. \frac{E_2(p)}{I_0(p)} \right|_{I_r=0}.$$

Thus $Z_{12}(p)$ represents the transfer impedance of N_1 when its output is open-circuited (i.e., $I_r = 0$). Since N_1 and N_2 are identical we also have, from $R_L = 1$ and reciprocity, $Z_{12}(p) = E_s/I_r$, where I_r is the current entering N_2 .

The impulse modulator is periodically operated every T seconds, and has the property that if its input is a continuous function $f(t)$ its output is a sequence of impulses:

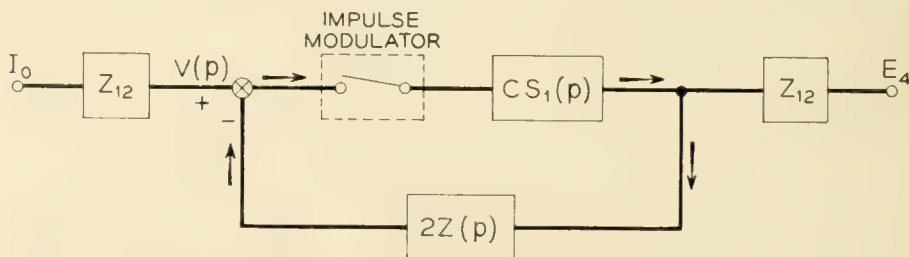
$$\sum_{-\infty}^{+\infty} f(t) \delta(t - kT).$$

The transfer function $S_1(p)$ is defined by

$$S_1(p) = \mathcal{L}[s_1(t)] = \frac{\omega_0^2}{p^2 + \omega_0^2} \cosh \frac{p\tau}{2} e^{-p\tau/2}. \quad (14)$$

Let $Z(p)$ be the driving point impedance at the terminal pair (2) of N_1 . It is also that of N_2 since N_1 and N_2 are assumed to be identical.

Let $V(p)$ be the output of the first block, then, by definition, $V(p) = Z_{12}(p)I_0$. Let $v(t)$ be the corresponding time function. The voltage $v(t)$



NOTE:
ALL BLOCKS ARE UNILATERAL

Fig. 4 — Zeroth-approximation block diagram.

is the output voltage of N_1 , when N_1 is excited by the current source I_0 and the switch S remains open at all times.

4.3 Analysis of the Block Diagram

For simplicity, suppose that the system starts from a relaxed condition (i.e., no energy stored) at $t = 0$. Let $z(t) = \mathcal{L}^{-1}[Z(p)]$. Considering the network N_1 as driven by i_0 and i_{r0} , it follows that the voltage $e_2(t)$ shown on Fig. 3 is given by

$$e_{20}(t) = v(t) - \int_0^t i_{r0}(t')z(t - t') dt'. \quad (15)$$

Similarly

$$e_{30}(t) = \int_0^t i_{r0}(t')z(t - t') dt'. \quad (16)$$

Thus

$$e_{20}(t) - e_{30}(t) = v(t) - 2 \int_0^t i_{r0}(t')z(t - t') dt'. \quad (17)$$

These equations have been derived by considering Fig. 1. They could have been also derived from the block diagram of Fig. 4 as follows: let $I_{r0}(p)$ be the output of $CS_1(p)$. As a result, the output of the block $2Z(p)$ is $2Z(p)I_{r0}(p)$. When this latter quantity is subtracted from $V(p)$ one gets $V(p) - 2Z(p)I_{r0}(p)$, which is the $\mathcal{L}$ -transform of the right-hand side of (17). Referring to the block diagram it is also seen that this quantity is the input to the impulse modulator.

Thus we see that if $I_{r0}(p)$ is the output of $CS_1(p)$, then the input of the impulse modulator is $e_{20}(t) - e_{30}(t)$ by virtue of (17). If this is the case the output of $CS_1(p)$ is given by $C[e_{20}(0) - e_{30}(0)]s_1(t)$, for $0 \leq t < T$, which, according to (9), is $i_{r0}(t)$.

Thus the block diagram of Fig. 4 is a convenient way of obtaining the zeroth approximation to the periodic steady-state solution.

In order to use the techniques developed for sampled data systems,^{1, 2} we introduce the following notation.² If $f(t) = \mathcal{L}^{-1}[F(p)]$, then we define $F^*(p)$ by the relation

$$F^*(p) = \frac{1}{T} \sum_{n=-\infty}^{+\infty} F(p + jn\omega_s), \quad (18)$$

where

$$\omega_s = \frac{2\pi}{T}. \quad (19)$$

If $f(0+)$ is defined by $\lim_{\epsilon \rightarrow 0} f(\epsilon^2)$, then, provided $f(0+) = 0$,*

$$F^*(p) = \mathcal{L} \left[\sum_{n=0}^{\infty} f(t) \delta(t - nT) \right]. \quad (20)$$

Going back to the system of Fig. 4 we get²

$$E_{40}(p) = \frac{[Z_{12}(p)I_0(p)]^*CS_1(p)Z_{12}(p)}{1 + 2C[S_1(p)Z(p)]^*}, \quad (21)$$

and

$$I_{ro}(p) = \frac{[Z_{12}(p)I_0(p)]^*CS_1(p)}{1 + 2C[S_1(p)Z(p)]^*}, \quad (22)$$

where according to the notation defined by (18)

$$[Z_{12}(p)I_0(p)]^* = \frac{1}{T} \sum_{n=-\infty}^{\infty} Z_{12}(p + jn\omega_s)I_0(p + jn\omega_s),$$

$$[S_1(p)Z(p)]^* = \frac{1}{T} \sum_{n=-\infty}^{+\infty} S_1(p + jn\omega_s)Z(p + jn\omega_s).$$

It should be stressed that (21) and (22) are not valid when τ is made identical to zero. When $\tau = 0$, $S_1(p) = 1$ for all p 's and since $Z(p) \sim 1/Cp$ as $p \rightarrow \infty$ the time function whose transform is $Z(p)S_1(p)$ is different from zero at $t = 0$. In such a case (20) does not hold. From a physical point of view, the feedback loop of Fig. 4 is unstable when τ is identically zero since an impulse generated by the impulse modulator produces instantaneously a step at the input of the impulse modulator. This step causes an instantaneous jump in the measure of the impulse at the output of the impulse modulator and so on. In short the feedback loop is unstable.

It should be pointed out that if the power density spectrum of I_0 is zero for frequencies higher than $\omega_s/2$, (21) reduces to

$$\frac{E_{40}}{I_0} = \frac{1}{T} \frac{CS_1(p)Z_{12}^2(p)}{1 + 2C[S_1(p)Z(p)]^*} \quad \text{for } |p| < \frac{\omega_s}{2}. \quad (23)$$

For certain applications it is convenient to rewrite (21) in a slightly

* When $f(0)$, as defined above, is different from zero, (20) should be replaced by

$$\mathcal{L} \left[\sum_{n=0}^{\infty} f(t) \delta(t - nT) \right] = F^*(p) + \frac{1}{2} f(0+).$$

different form. Advancing the time function $s_1(t)$ by $\tau/2$ seconds, one gets the function $s_0(t)$ which is even in t . As a result its transform $S_0(p)$ is purely real, that is,

$$S_0(p) = \frac{\omega_0^2}{p^2 + \omega_0^2} \cosh \frac{p\tau}{2}.$$

From an analysis carried out in detail in Appendix IV we finally obtain

$$E_{40}(p) = \frac{[Z_{12}(p)I_0(p)]^* S_0(p) Z_{12}(p)}{2[Z(p)S_0(p)]^*}. \quad (24)$$

It should be pointed out that (23) is still valid when $\tau = 0$. Equations (20) and (23) give the zeroth approximation to the gain of the system for any driving current $i_0(t)$.

In many cases it is sufficient to know only the steady-state response $E_{40}(p)$ to an input $i_0(t) = I_0 e^{j\omega_0 t}$. The response $E_{40}(p)$, as given by (23) [or (20)] includes both transient and steady-state terms. Since $I_0(p) = \frac{I_0}{p - j\omega_0}$ equation (24) gives

$$E_{40}(p) = \frac{\left(\frac{1}{T} \sum_{n=-\infty}^{+\infty} Z_{12}(p + jn\omega_s) \frac{I_0}{p + jn\omega_s - j\omega_0} \right) S_0(p) Z_{12}(p)}{2[Z(p)S_0(p)]^*}. \quad (25)$$

Since neither $S_0(p)$ nor $Z_{12}(p)$ have poles on the imaginary axis, the steady state includes only the terms corresponding to the imaginary axis poles of the summation terms. Thus the steady-state response is of the form

$$\sum_{n=-\infty}^{+\infty} A_n e^{j(\omega_0 - n\omega_s)t},$$

where, from (25),

$$A_n = \frac{I_0 Z_{12}(j\omega_0) S_0(j\omega_0 - jn\omega_s) Z_{12}(j\omega_0 - jn\omega_s)}{2 \sum_{k=-\infty}^{+\infty} Z[j\omega_0 + j(k - n)\omega_s] S_0[j\omega_0 + j(k - n)\omega_s]}. \quad (26)$$

V. TRANSMISSION LOSS

A practically important question is to find out *a priori* whether a switched filter necessarily introduces some transmission loss.

The following considerations apply exclusively to the zeroth order approximation. It will be shown that assuming ideal elements, the transmission at dc may have as small a loss as desired.

By transmission at dc we mean the ratio of the dc component of the steady state output voltage to the intensity of the applied direct current. Thus we refer to (26) and set $\omega_0 = 0$ and $n = 0$. Suppose the lossless networks N_1 and N_2 are designed so that their transfer impedance Z_{12} is of the Butterworth type, that is

$$|Z_{12}(j\omega)|^2 = \frac{1}{1 + \omega^{2M}},$$

where for our purposes M is a large integer.

In the following sum, which is the denominator of (26) when $\omega_0 = n = 0$

$$2 \sum_{k=-\infty}^{+\infty} Z(jk\omega_s)S_0(jk\omega_s),$$

(where $\omega_s > 2$ since the cutoff of the networks N occurs at $\omega = 1$), the terms corresponding to values of $k \neq 0$ will make a contribution that vanishes as $M \rightarrow \infty$. This is a consequence of the following facts:

(a) $\text{Re}[Z(jk\omega_s)] = |Z_{12}(jk\omega_s)|^2$, since the networks N_1 and N_2 are dissipationless. Hence for $k \neq 0$ and as $M \rightarrow \infty$ $\text{Re}[Z(jk\omega_s)] \rightarrow 0$,

(b) $\text{Im}[Z(jk\omega_s)] = -\text{Im}[Z(-jk\omega_s)]$,

(c) $S_0(j\omega)$ is real.

Thus the imaginary part of the products $Z(jk\omega_s)S_0(jk\omega_s)$ cancel out and the real part (for $k \neq 0$) decreases exponentially to zero as $M \rightarrow \infty$. Hence for sufficiently large M the denominator of (26) may be made as close to two as desired.

It is easy to check that the numerator of (26) reduces to I_0 , the intensity of the applied direct current. Therefore the ratio of A_0 , the dc component of the output voltage to I_0 may be made as close to one-half as desired.

VI. A SIMPLE EXAMPLE

Since the approximate formulae derived in Section IV are somewhat unfamiliar it seems proper to consider in a rather detailed manner a simple example.*

Consider the system of Fig. 5. Assume that the current source applies a constant current to the system and assume that the steady state is reached. For simplicity let $I_0 R = E$.

The steady-state behavior of the voltages $e_2(t)$ and $e_3(t) = e_4(t)$ is

* In addition, the limiting case of the sampling rate $\rightarrow \infty$, i.e., $T \rightarrow 0$, is treated in Appendix II.

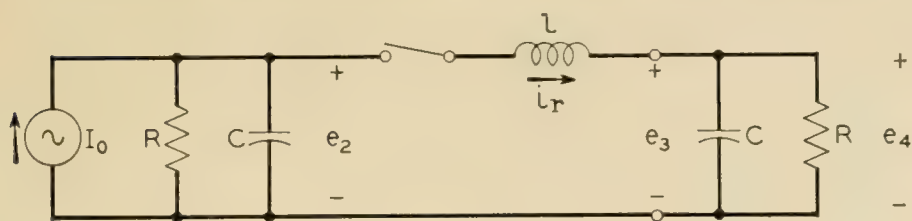


Fig. 5 — A simple example.

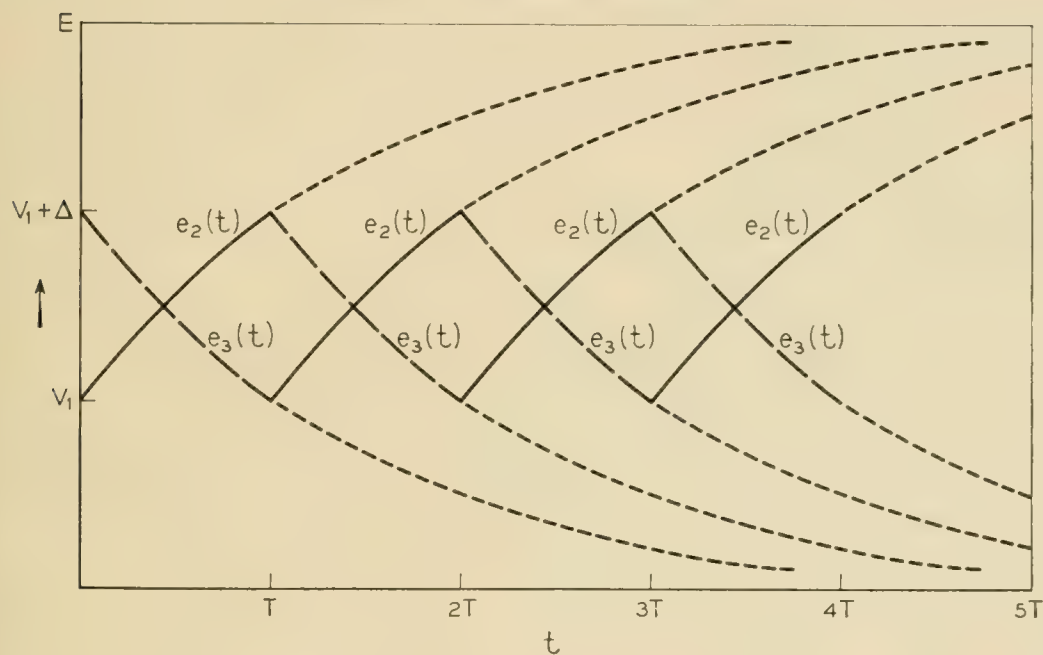


Fig. 6 — Waveforms of the network of Fig. 5.

shown on Fig. 6. It is further assumed that the duration τ during which the switch is closed is negligible compared to T , the interval between two successive closures.

Let $\bar{e}_4$ be the average value of the steady-state voltage $e_4(t)$. Thus $\bar{e}_4$ is equal to A_0 as given by (26) with $\omega_0 = n = 0$, namely,

$$e_4 = I_0 \frac{Z_{12}^2(0) S_0(0)}{2 \sum_{-\infty}^{+\infty} Z(jk\omega_s) S_0(jk\omega_s)} .$$

In this particular case

$$Z(p) = Z_{12}(p) = \frac{R}{1 + pRC} = \frac{C^{-1}}{p + \frac{1}{RC}} .$$

Since we assume τ to be infinitesimal $S_0(p)$ and $S_1(p)$ may be considered equal to unity over the band of interest. Using the expansion

$$\coth z = \frac{1}{z} + \sum_1^{\infty} \frac{2z}{z^2 + n^2\pi^2},$$

we obtain

$$Z^*(p) = \frac{1}{T} \sum_{-\infty}^{+\infty} \frac{C^{-1}}{(p + jn\omega_s) + \frac{1}{RC}} = \frac{1}{2C} \coth \left[\left(p + \frac{1}{RC} \right) \frac{T}{2} \right].$$

Hence

$$\frac{2}{T} \sum_{-\infty}^{+\infty} Z(jk\omega_s) = 2Z^*(p) \Big|_{p=0} = \frac{1}{C} \coth \left(\frac{T}{2RC} \right).$$

Thus finally

$$\bar{e}_4 = \frac{I_0 R^2}{T \frac{1}{C} \coth \left(\frac{T}{2RC} \right)} = E \frac{RC}{T} \frac{1}{\coth \left(\frac{T}{2RC} \right)}. \quad (27)$$

This last result obtained from the theory developed above is now going to be checked directly. Referring to Fig. 6, where the notation is defined, and noting the periodicity of the boundary conditions, we get

$$\begin{aligned} (V_1 + \Delta)e^{-(T/RC)} &= V_1, \\ (E - V_1)e^{-(T/RC)} &= E - (V_1 + \Delta). \end{aligned}$$

Noting that $e_4(t) = (V_1 + \Delta)e^{-t/RC}$, and solving for V_1 and Δ we finally get

$$e_4(t) = \frac{e^{-t/RC}}{1 + e^{-T/RC}}.$$

By definition

$$\bar{e}_4 = \frac{1}{T} \int_0^T e_4(t) dt = \frac{E RC}{T} \frac{1 - \bar{e}^{-T/RC}}{1 + \bar{e}^{-T/RC}},$$

or

$$\bar{e}_4 = E \frac{RC}{T} \frac{1}{\coth \left(\frac{T}{2RC} \right)}.$$

This last equation checks with (27).

VII. NUMERICAL EXAMPLES

Consider the network of Fig. 7. The cutoff of both N_1 and N_2 occurs at $\omega = 1$. In view of the sampling theorem good transmission requires

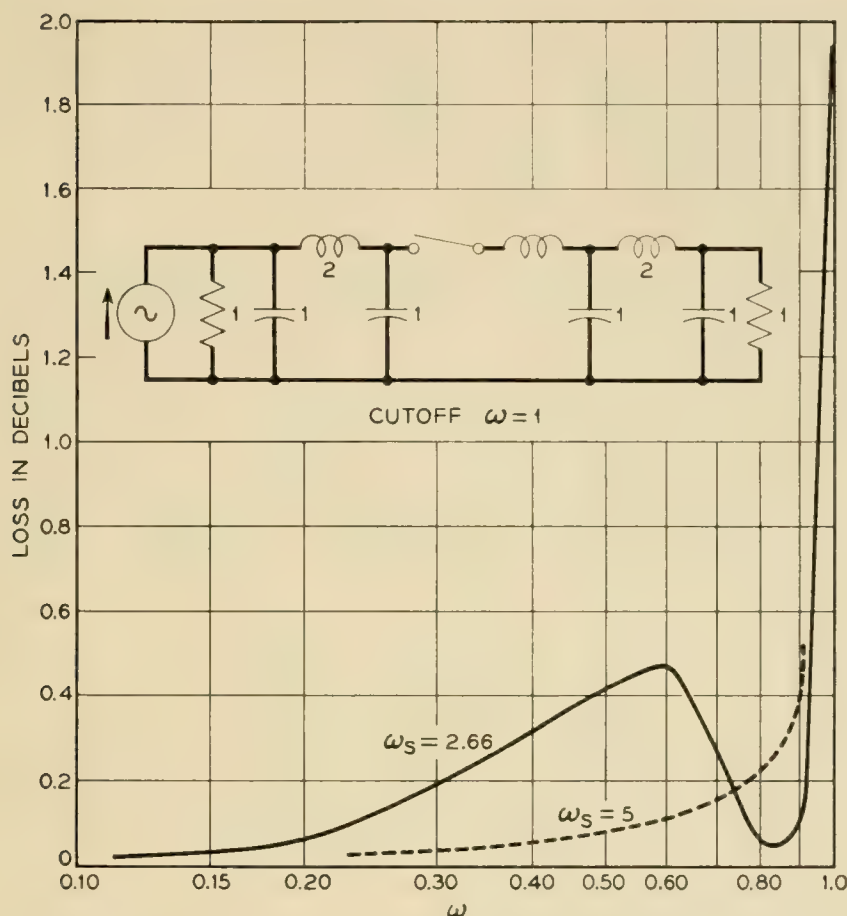


Fig. 7 — Computed transmission loss.

that the signal be sampled at a rate at least twice as large as its highest frequency component. Since the cutoff occurs at $\omega = 1$, the sampling angular frequency should at least be equal to 2. For illustration purposes we have taken $\omega_s = 2.67$ and $\omega_s = 5$ for the angular sampling frequency. The value $\omega_s = 2.67$ corresponds to a cutoff at 3 kc and a sampling rate of 8 kc. The ratio τ/T was taken to be $1/125$. The transmission through the switched network as given by the zeroth approximation is shown for both cases on Fig. 7.

As expected the transmittance of the switched filter gets closer to that of an ordinary filter as the switching frequency increases.

VIII. THE SUCCESSIVE APPROXIMATION SCHEME

The ideas involved in the successive approximation scheme are simple and straightforward. One point remains to be settled, namely the convergence of the procedure.

We shall assign a subscript 1 to the correction to be applied to the

zeroth approximation in order to obtain the first approximation. Thus adding $i_{r1}(t)$ to $i_{r0}(t)$ we get the first approximation $i_{r0}(t) + i_{r1}(t)$. More generally the k th approximation is $\sum_{n=0}^k i_{rn}(t)$. The procedure will converge if, in particular, the infinite series $\sum_{n=0}^{\infty} i_{rn}(t)$ converges.

8.1 Preliminary Steps

(a) Let us normalize the frequency (and consequently the time) so that the switching period T is unity. Since the networks N_1 and N_2 must have high insertion loss for $\omega > \frac{1}{2}(2\pi/T) = \pi$, the pass band of N_1 and N_2 must be the order of 1 radian/sec. As a result the element values of the capacitor C and the inductance L_n (see Fig. 3) are also 0(1).

(b) For the excitation $i_0 = e^{i\omega t}$, the zeroth approximation derived above may be written in terms of Fourier components:

$$i_{r0}(t) = e^{i\omega t} \sum_{k=-\infty}^{+\infty} I_{r0,k} e^{i2\pi k t},$$

$$i_{n0}(t) = e^{i\omega t} \sum_{k=-\infty}^{+\infty} I_{n0,k} e^{i2\pi k t}.$$

Let $\bar{i}_{r0}$ denote the complex conjugate of i_{r0} , then

$$i_{r0}(t) \bar{i}_{r0}(t) = \sum_{k, \ell=-\infty}^{+\infty} I_{r0,k} \bar{I}_{r0,\ell} e^{2\pi i(k-\ell)t}.$$

Since the functions $e^{2\pi i k t}$ [$k = \dots -1, 0, 1, \dots$] are orthonormal over the interval $(0, 1)$ and form a complete set,³ we have from Bessel's equality:³

$$\int_0^1 |i_{r0}(t)|^2 dt = \sum_{k=-\infty}^{+\infty} |I_{r0,k}|^2 = N(I_{r0}),$$

where $N(I_{r0})$ denotes the norm of the vector I_{r0} which is defined by its components $I_{r0,k}$ ($k = \dots -1, 0, +1 \dots$). Similarly,

$$\int_0^1 |i_{n0}(t)|^2 dt = \sum_{k=-\infty}^{+\infty} |I_{n0,k}|^2 = N(I_{n0}).$$

(c) Since the switch is periodically closed we shall be interested in the Fourier series expansion of $\Delta(t)$:

$$\Delta(t) = u(t) - u(t - \tau) = \sum_{-\infty}^{+\infty} \Delta_k e^{ik2\pi t},$$

where $\Delta_0 = \tau$ and

$$\sum_{-\infty}^{+\infty} |\Delta_k|^2 = \tau.$$

Since $\tau/T \ll 1$, and since the frequency has been normalized so that $T = 1$, we have $\tau \ll 1$.

Using the convolution in the frequency domain, we have

$$i_{n0}(t)\Delta(t) = \sum_{k=-\infty}^{+\infty} \left(\sum_{\alpha=-\infty}^{+\infty} \Delta_{k-\alpha} I_{n0,\alpha} \right) e^{i(\omega+2\pi k)t}.$$

If we introduce the infinite matrix G defined by

$$G_{ik} = \Delta_{i-k} \quad (i, k = -\infty, \dots, -1, 0, +1, \dots, \infty),$$

the convolution may be represented by the product, GI_{n0} , where I_{n0} is the vector whose components are $I_{n0,k}$ ($k = \dots, -1, 0, 1, \dots$).

(d) Considering the network shown on Fig. 3, let $E(p)$ be the ratio of $I_n'(p)$ to $I_r(p)$. Taking into account the assumed identity between N_1 and N_2 it follows that

$$\left. \frac{+I_n(p)}{I_r(p)} \right|_{I_0=0} = \frac{I_n'(p)}{I_r(p)} = E(p).$$

Using the system of (1) and, for example, by Neumann series expansion of the inverse matrix, we get

$$E(p) = \frac{1}{L_n C p^2} - \frac{2}{L_n^2 C^2 p^4} + \dots$$

(e) Considering now the effect of $i_n(t)$ and $i_n'(t)$ on $i_r(t)$, (42) of Appendix I gives $I_r(p)$ as a function of $I_n(p)$ and $I_n'(p)$. In the present discussion where we are interested in the steady state of $i_r(t)$ it is essential to keep in mind that since the switch opens at $t = \tau$, the memory of the resonant circuit extends only over an interval $0 < t \leq \tau$. To take this into account we must modify the factor $(\omega_0^2/2)/(p^2 + \omega_0^2)$ of (40), because the impulse response (which represents this memory) must be identically zero for $t > \tau$. The resulting new expression is

$$F(p) = \frac{\frac{\omega_0^2}{2}}{p^2 + \omega_0^2} e^{-p\tau/2} [e^{p\tau/2} + e^{-p\tau/2}],$$

or

$$F(p) = \frac{\omega_0^2}{p^2 + \omega_0^2} e^{-p\tau/2} \cosh \frac{p\tau}{2}.$$

Since the time function whose transform is $F(p)$ is non-negative for all t 's and since $F(0) = 1$, it follows that

$$|F(j\omega)| \leq 1. \quad (28)$$

8.2 *Matrix description of the successive approximations*

From the developments of Section IV, we know $i_{r0}(t)$, $i_{n0}(t)$ and $i_{n0}'(t)$ or what is equivalent, the vectors I_{r0} , I_{n0} and I_{n0}' . The first approximation takes into account the effect of $i_{n0}(t)$ and $i_{n0}'(t)$ on $i_r(t)$. [See equation (1.c) and (1.d)]. The time functions $i_{n0}(t)$ and $i_{n0}'(t)$ affect the system R only during the interval $(0, \tau)$. Therefore we must consider the vector $G(I_{n0} + I_{n0}')$ which corresponds to the excitation of the resonant circuit. Since the opening of the switch after a closure time τ forcibly brings $i_r(t)$ to zero we have

$$I_{r1} = G F G (I_{n0} + I_{n0}'), \quad (29)$$

where the matrix G has been defined above and the matrix F is a diagonal matrix whose diagonal elements F_k ($k = \dots -1, 0, +1 \dots$) are defined by $F_k = F(j\omega_s + j2\pi k)$. Note that (28) implies that $|F_k| \leq 1$ for all k 's. It should be kept in mind that $I_{r0} + I_{r1}$ is the first approximation to the exact $I_r(p)$.

The next iteration is obtained by first taking into account the effect of I_{r1} on the rest of the network:

$$\begin{aligned} I_{n1} &= E I_{r1}, \\ I_{n1}' &= E I_{r1}, \end{aligned} \quad (30)$$

where E is a diagonal matrix whose elements E_k ($k = \dots, -1, 0, +1, \dots$) are defined by $E_k = E(j\omega_s + 2\pi k j)$, and then the effects of I_{n2} and I_{n2}' on I_r , that is,

$$I_{r2} = G F G (I_{n1} + I_{n1}'), \quad (31)$$

combining (30) and (31), $I_{r2} = 2 G F G E I_{r1}$. A repetition of the same procedure would lead to $I_{r3} = 2 G F G E I_{r2}$, and in general $I_{rn+1} = 2 G F G E I_{rn}$.

Since the n th approximation to $I_r(p)$ is given by the sum $\sum_{k=0}^n I_{rk}$, the successive approximation scheme will be convergent only if the series

$$\sum_{k=0}^{\infty} I_{rk}$$

converges. This will be the case if and only if the series

$$[1 + 2 G F G E + \dots + (2 G F G E)^n + \dots] I_{r1} \quad (32)$$

converges.

8.3 Convergence Proof

Consider a vector X of bounded norm corresponding to a time function $x(t)$ having the property that $x(t) = 0$ for $\tau \leq t \leq T$ and $x(t) \neq 0$ for $0 < t < \tau$. In the above scheme, the vector X would be I_{rn} . Let us define the vectors Y , Z , U and V by the relations

$$Y = EX, \quad (33)$$

$$Z = GY, \quad (34)$$

$$U = FZ, \quad (35)$$

$$V = 2GU, \quad (36)$$

hence

$$V = 2GFGE X. \quad (37)$$

We wish to show that $N(V) \leq \alpha N(X)$ with $\alpha < 1$, since these inequalities imply that the infinite series (32) converges.

Since (a) N_1 and N_2 are low-pass filters with cutoff $\leq \pi$ radians/sec, (b) $E(p) = 1$ for $p = 0$, (c) $E(p) \propto 1/L_n C p^2$ for $p \gg 1$, only a few of the E_k 's will be of the order of unity. In most cases E_{-1} , E_0 , E_1 will be smaller than unity, thus,

$$N(Y) \leq N(X). \quad (38)$$

In view of the pulsating character of $x(t)$ the power spectrum of $x(t)$ is almost constant up to frequencies of the order of π/τ radians/sec. Because of the low-pass characteristic of $E(p)$, the function $y(t)$ associated with the vector Y is smooth in comparison to $x(t)$, thus from (34),

$$N(Z) = \int_0^1 |z(t)|^2 dt = \int_0^\tau |y(t)|^2 dt = a\tau N(Y),$$

where $a = 0(1)$.

Since $|F_k| \leq 1$ for all k 's, from (33), $N(U) \leq N(Z)$, hence $N(U) = b\tau N(Y)$ with $b = 0(1)$.

$$N(U) = b\tau N(Y) \quad \text{with} \quad b = 0(1).$$

From (36) we have

$$N(V) = 2 \int_0^\tau |u(t)|^2 dt \leq 2 \int_0^1 |u(t)|^2 dt = 2N(U).$$

Thus we finally get

$$N(V) = 2b\tau N(Y) \quad \text{where} \quad b = 0(1), \quad (39)$$

and since $\tau \ll 1$ we get from (38) and (39) $N(V) = \alpha N(X)$ with $\alpha < 1$. Hence the convergence is established.

IX. A MODIFICATION OF THE BLOCK DIAGRAM TO IMPROVE THE ZEROth APPROXIMATION

In principle it is possible to obtain a block diagram whose transmission characteristic is equal to the first approximation. In many cases it is not necessary to go that far. The first approximation takes into account the effect of the currents $i_{n0}(t)$ and $i_{n0}'(t)$ on the resonant circuit of Fig. 2. Since during the switch closure time the currents i_{n0} and i_{n0}' cannot vary much, let us assume that they remain constant for the duration of the switch closure.

Referring to the analysis of Appendix I and to (42) in particular, we see that the current i_r is increased by

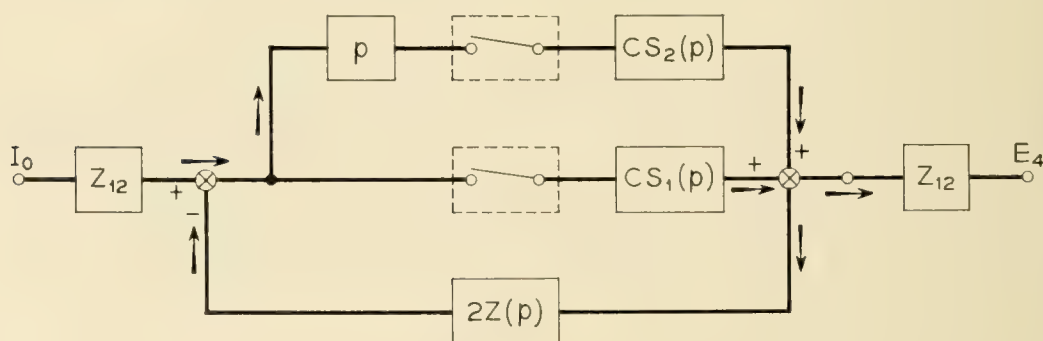
$$\delta i_n(p) = \frac{\omega_0^2}{p^2 + \omega_0^2} \frac{i_n(0-) + i_n'(0-)}{2p},$$

or

$$\delta i(t) = C \frac{e_2(0-) - e_3(0-)}{2} (1 - \cos \omega_0 t) \quad 0 \leq t < \tau.$$

Defining $S_2(p) = \mathcal{L}^{-1}\{\frac{1}{2}(1 - \cos \omega_0 t)[u(t) - u(t - \tau)]\}$, or

$$S_2(p) = \frac{\omega_0^2}{2} \frac{1}{p(p^2 + \omega_0^2)} - \frac{(2p^2 + \omega_0^2)}{2p(p^2 + \omega_0^2)} e^{-p\tau},$$



$$E_4(p) = \frac{C[Z_{12}(p)I_0(p)]^* S_1(p) + C[pZ_{12}(p)I_0(p)]^* S_2(p)}{1 + 2C\left\{[Z(p)S_1(p)]^* + [pZ(p)S_2(p)]^*\right\}}$$

Fig. 8 — Modified block diagram.

and recalling that the input of the impulse modulator of Fig. 4 is $e_2(0) - e_3(0)$, it becomes obvious that the modified block diagram should be that given by Fig. 8. The output of the modified block diagram is given by²

$$E_4(p) = \frac{C[Z_{12}(p)I_0(p)]^*S_1(p) + C[pZ_{12}(p)I_0(p)]^*S_2(p)}{1 + 2C\{[Z(p)S_1(p)]^* + [pZ(p)S_2(p)]^*\}}.$$

X. CONCLUSION

Let us compare the method of solution presented above with the more formal approach proposed by Bennett. The latter method leads to the exact steady-state transmission through a network containing periodically operated switches. This method is perfectly general in that it does not require any assumption relative to the properties of the network nor to the ratio of τ/T . As expected this generality implies a lot of detailed computations. In particular it requires, for each reactance of the network, the computation of the voltage across it due to any initial condition. The method presented in this paper is not so general because it assumes first that the ratio τ/T is small; second the value of the inductance ℓ is very much smaller than that of L_n (see Fig. 3). The result of these assumptions is that the system of time varying equations may be solved by successive approximations with the further advantage that the convergence proof guarantees that, for very small τ/T , the zeroth approximation will be a close estimate of the exact solution.

The zeroth approximation may conveniently be obtained by considering a block-diagram analogous to those used in the analysis of sampled servomechanisms. Further the proposed method leads directly to some interesting results, for example, as far as the zeroth approximation is concerned, the dc transmission may be achieved with as small a loss as desired provided the lossless networks N_1 and N_2 are suitably designed. Another advantage of the proposed method is that the simplicity of the analysis permits the designer to investigate at a small cost a large number of possible designs.

Finally it should be pointed out that this approach to the solution of a system of time-varying linear differential equations may find applications in many other physical problems.

APPENDIX I

ANALYSIS OF THE RESONANT CIRCUIT

Consider the resonant circuit of Fig. 2. Suppose that at $t = 0$, the left-hand capacitor has a potential $e_2(0)$ and the right-hand capacitor has

a potential $e_3(0)$ and that at $t = 0$ the current i_r through the inductance ℓ is zero.

The network equation is

$$\ell \frac{d}{dt} i_r + \frac{2}{C} \int i_r dt = 0.$$

Now $i_r(0) = 0$ and $di_r(0)/dt = [e_2(0) - e_3(0)]/\ell$. Let $2/\ell C = \omega_0^2$, then $di_r(0)/dt = \omega_0^2 C [e_2(0) - e_3(0)]/2$.

Using Laplace transforms,

$$(p^2 + \omega_0^2)I_r(p) = pi_r(0) + \frac{di_r(0)}{dt}, \quad (40)$$

$$I_r(p) = \frac{\omega_0^2 C}{2} [e_2(0) - e_3(0)] \frac{1}{p^2 + \omega_0^2},$$

hence

$$i_r(t) = \omega_0 C \frac{e_2(0) - e_3(0)}{2} \sin \omega_0 t \quad (41)$$

and

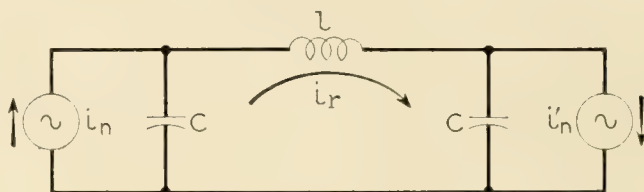
$$q(t) = \int_0^t i_r(t) dt = C \frac{e_2(0) - e_3(0)}{2} [1 - \cos \omega_0 t].$$

If $2\pi/\omega_0 = 2\tau$, i.e., $\tau = \pi\sqrt{\ell C/2}$, which means that the duration of the switch closure is a half-period of the resonance of the tuned circuit, then

$$i_r(t) = \frac{\pi}{\tau} \frac{C[e_2(0) - e_3(0)]}{2} \sin \frac{\pi t}{\tau},$$

$$q(t) = \frac{C[e_2(0) - e_3(0)]}{2} \left[1 - \cos \frac{\pi t}{\tau} \right].$$

It is clear then that, during the period τ , the charge transferred onto the



$$I_r(p) = \frac{\omega_0^2}{p^2 + \omega_0^2} \frac{I_n(p) + I_n'(p)}{2}$$

Fig. 9 — Resonant circuit excited by current sources I_n and I_n' .

right-hand capacitor is $q(\tau) = C[e_2(0) - e_3(0)]$ and as a result at time $t = \tau$ the right-hand capacitor has a voltage $e_2(0)$ and the left-hand capacitor has a voltage $e_3(0)$. Considering now the network of Fig. 9, the equation is

$$\ell \frac{d^2 i_r}{dt^2} + \frac{2}{C} i_r = \frac{1}{C} [i_n(t) + i_n'(t)].$$

Assuming all initial conditions* to be zero we get,

$$I_r(p) = \frac{\omega_0^2}{p^2 + \omega_0^2} \frac{I_n(p) + I_n'(p)}{2}. \quad (42)$$

APPENDIX II

STUDY OF THE LIMITING CASE $T \rightarrow 0$

We expect that if the sampling period $T \rightarrow 0$, which is equivalent to stating that the sampling frequency $\omega_s \rightarrow \infty$, then the inductance $\ell \rightarrow 0$ and as a result the voltage $e_3(t)$ will be infinitely close, at all times, to the voltage $e_2(t)$. Thus, in the limit, everything happens as if the terminal pairs (2) of N_1 and N_2 were directly connected. In that case the gain of the system is

$$\frac{Z_{12}(p)}{2Z(p)} Z_{12}(p),$$

as is easily seen by referring to the Thevenin equivalent circuit of N_1 .

Let us show that as $T \rightarrow 0$, (21) leads to the same result. First note that both $Z_{12}I_0$ and ZS_1 go to zero at least as fast as $1/p^2$ for $p \rightarrow \infty$. Hence the summations in (21) reduce to the term corresponding to $n = 0$. Therefore,

$$E_4(p) = \frac{C[Z_{12}I_0]^*S_1}{1 + 2C[ZS_1]^*} Z_{12} \rightarrow \frac{CZ_{12}I_0Z_{12}S_1}{T + 2CZS_1} \rightarrow \frac{Z_{12}^2}{2Z} I_0.$$

APPENDIX III

ZEROth APPROXIMATION IN THE CASE WHERE N_1 IS NOT IDENTICAL TO N_2

Let, for $k = 1, 2$; C_k be the shunt capacitor at the terminal pair 2 of N_k , $Z_k(p)$ be the driving point impedance of N_k , and $Z_{12}^{(k)}(p)$ be the transfer impedance of N_k . In the present case the capacitors C_1 and C_2 are in series in the resonant circuit of Fig. 2. It can be shown that the

* Their contribution has been found in (40).

charge exchanged during one-half period of the resonance is

$$\frac{2C_1C_2}{C_1 + C_2} [e_2(0) - e_3(0)].$$

For the present case, (16) and (17) become

$$e_2(t) = v(t) - \int_{-\infty}^t i_{r0}(\tau) z_1(t - \tau) d\tau,$$

$$e_3(t) = \int_{-\infty}^t i_{r0}(\tau) z_2(t - \tau) d\tau.$$

Following the same procedure as before we are finally led to the block diagram of Fig. 10 whose output is given by

$$E_4(p) = \frac{[Z_{12}^{(1)}(p)I_0(p)]^* \frac{2C_1C_2}{C_1 + C_2} S_1(p)Z_{12}^{(2)}(p)}{1 + \frac{2C_1C_2}{C_1 + C_2} \{[Z_1(p) + Z_2(p)]S_1(p)\}^*}.$$

APPENDIX IV

THE DERIVATION OF EQUATION (24)

Considering the method used in Section IV to derive the zeroth approximation, it is clear that during the switch closure the voltages $e_2(t)$ and $e_3(t)$ vary sinusoidally, that is,

$$e_2(t) = e_2(0) - \frac{e_2(0) - e_3(0)}{2} \left[1 - \cos \frac{\pi t}{\tau} \right],$$

$$e_3(t) = e_3(0) + \frac{e_2(0) - e_3(0)}{2} \left[1 - \cos \frac{\pi t}{\tau} \right].$$

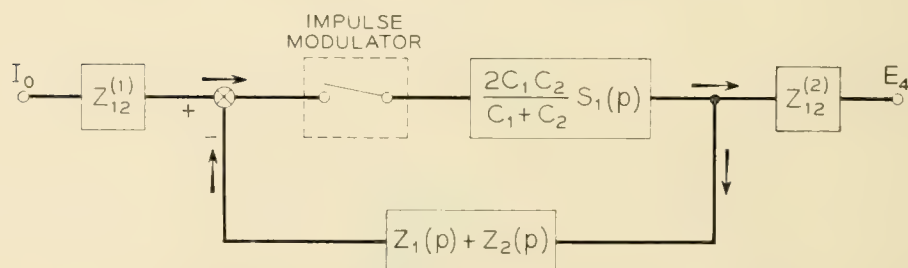


Fig. 10 — Zeroth approximation: modified block diagram for the case where N_1 and N_2 are not identical.

Thus, it always happens that for $t = \tau/2$, i.e., at the middle of switch closure time, $e_2(t) - e_3(t) = 0$.

Therefore if we consider the time function $e_2(t) - e_3(t)$ we have for all k 's $(-\infty, \dots 0, \dots +\infty)$,

$$e_2\left(kT + \frac{\tau}{2}\right) - e_3\left(kT + \frac{\tau}{2}\right) = 0.$$

If, for simplicity of analysis, we assume that the switch is closed during the intervals $-(\tau/2) + kT \leq t \leq +\tau/2 + kT$, then for all k 's,

$$e_2(kT) - e_3(kT) = 0.$$

Using (17), this condition implies that $[V(p)]^* - 2[I_{r0}(p)Z(p)]^* = 0$.

Now, remembering that $i_{r0}(t)$ consists of a sequence of half sine waves whose shape is defined by $s_0(t)$ (which is by definition identical to $s_1(t)$ except for an advance in time of $\tau/2$) it follows that $I_{r0}(p) = B(p)S_0(p)$, where $B(p)$ is the $\mathcal{L}$ -transform of the sequence of impulses whose measure is equal to the charge interchanged between N_1 and N_2 at each switch closure. Since $[B(p)S_0(p)Z(p)]^* = B(p)[S_0(p)Z(p)]^*$, then

$$B(p) = \frac{[Z_{12}(p)I_0(p)]^*}{2[S_0(p)Z(p)]^*}.$$

From which it immediately follows that

$$I_{r0}(p) = \frac{[Z_{12}(p)I_0(p)]^*S_0(p)}{2[S_0(p)Z(p)]^*}$$

and

$$E_{40}(p) = \frac{[Z_{12}(p)I_0(p)]^*S_0(p)Z_{12}(p)}{2[S_0(p)Z(p)]^*},$$

where

$$[S_0(p)Z(p)]^* = \frac{1}{T} \sum_{n=-\infty}^{+\infty} S_0(p + jn\omega_s) Z(p + jn\omega_s).$$

REFERENCES

1. L. A. MacColl, *Fundamental Theory of Servomechanisms*, D. Van Nostrand Co. Inc., New York, 1945.

2. John G. Truxal, *Automatic Feedback Control System Synthesis*, McGraw-Hill Book Co. Inc. New York, 1955.
3. R. Courant and D. Hilbert, *Methods of Mathematical Physics*, Interscience Publishers, Inc. New York, 1953.
4. W. R. Bennett, Steady-state Transmission Through Networks Containing Periodically Operated Switches, I.R.E. Trans., **PGCT-2**, No. 1, pp. 17-22, March, 1955.
5. T. R. Bashkow, The A-matrix, New Network Description, to be published in I.R.E. Trans., **PGCT-4**, No. 3, 1957.

Experimental Transversal Equalizer for TD-2 Radio Relay System

By B. C. BELLOWS and R. S. GRAHAM

(Manuscript received February 26, 1957)

To determine the effect of improved equalization on the performance of TD-2 radio relay systems, an experimental adjustable transversal equalizer has been developed. The equalizer is based on the echo principle as used in transversal filters, and operates in the 60- to 80-mc frequency band. Seven pairs of adjustable leading and lagging echo terms provide flexibility for simultaneous gain and delay equalization. Directional couplers are used for tapping and controlling the echo voltages. Field experiments have shown that system equalization can be improved appreciably by the use of such equalizers.

INTRODUCTION

The TD-2 radio relay system¹ employs frequency modulation to transmit multichannel telephony or television. The frequency modulated signal requires a transmission system whose gain and envelope delay are constant over the frequency band, 20 mc wide, used to transmit the signal. Deviations from constant gain or delay result in non-linear distortion of the demodulated signal. For television this results in distorted images, and for multichannel telephone transmission this introduces cross modulation among the voice channels.

Basic equalization is provided in each repeater. In addition, certain fixed equalizers have been used on a mop-up basis. However, there remains some residual distortion of random shape. This paper discusses the design of an adjustable equalizer to compensate for this distortion.

I. BASIC EQUALIZATION

The transmission path through a TD-2 repeater consists of an intermediate frequency portion covering the 60- to 80-mc band and radio frequency channels 20-mc wide in the 3,700- to 4,200-mc band. At each

repeater the RF channels are separated by wave guide filters and each is demodulated to the IF frequency for amplification and equalization. The outgoing signal is then modulated back to an RF channel, at a different frequency from the incoming signal to reduce interference. The RF and IF portions of the repeater were designed so that the combination would have flat gain to within the closest practicable limits over the 20-mc band, and a number of adjustments are provided in the IF amplifier to maintain this flatness under field conditions. The unequalized repeater, however, has an envelope delay distortion characteristic shown in Fig. 1 which is approximately parabolic. To minimize this distortion, each repeater contains a 315A equalizer, which has approximately the inverse of the delay distortion of a typical repeater.

II. SUPPLEMENTARY EQUALIZATION

In a TD-2 system consisting of many repeaters in tandem, both gain and delay distortions may accumulate to the point where additional equalization is necessary. If the pass band of the repeaters shifts slightly in frequency, due to changes in temperature or adjustment, the difference between the repeater delay and the delay of its equalizer will result in delay distortion which has approximately a linear slope with frequency. This may be corrected at main repeater stations by combina-

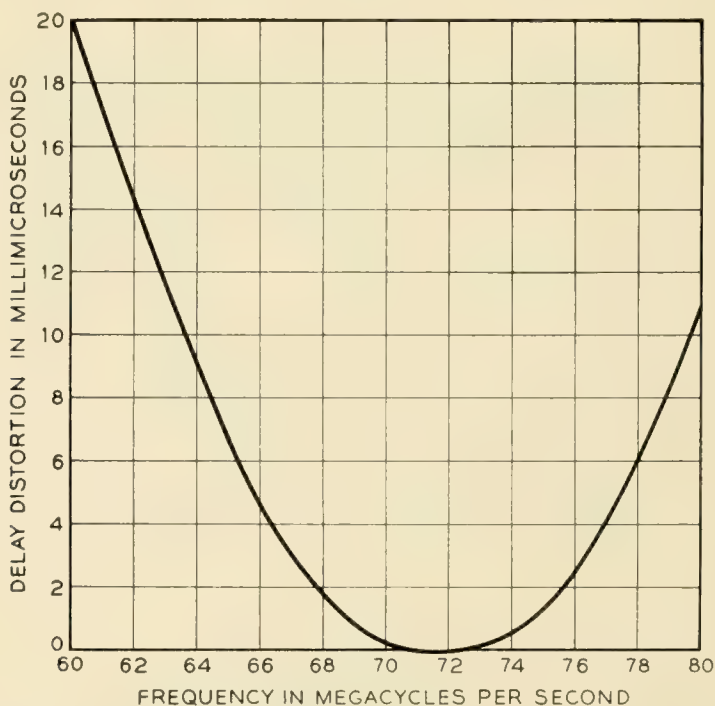


Fig. 1 — Over-all delay distortion of a typical microwave repeater, TD-2 system.

tions of delay slope equalizers. The characteristics of the 319A, B and C equalizers provided for this purpose are shown in Fig. 2. Each equalizer consists of two bridged-T all-pass sections. There is also some variation in bandwidth of the TD-2 repeaters, resulting in part from the fact that the waveguide filters used in the higher frequency radio channels are somewhat broader than those in the lower frequency channels. This variation in bandwidth results in delay distortion which has a parabolic shape with frequency. Use of a larger or smaller number of the basic 315A equalizers corrects this.

Over long circuits, small distortions in the gain shape of the TD-2

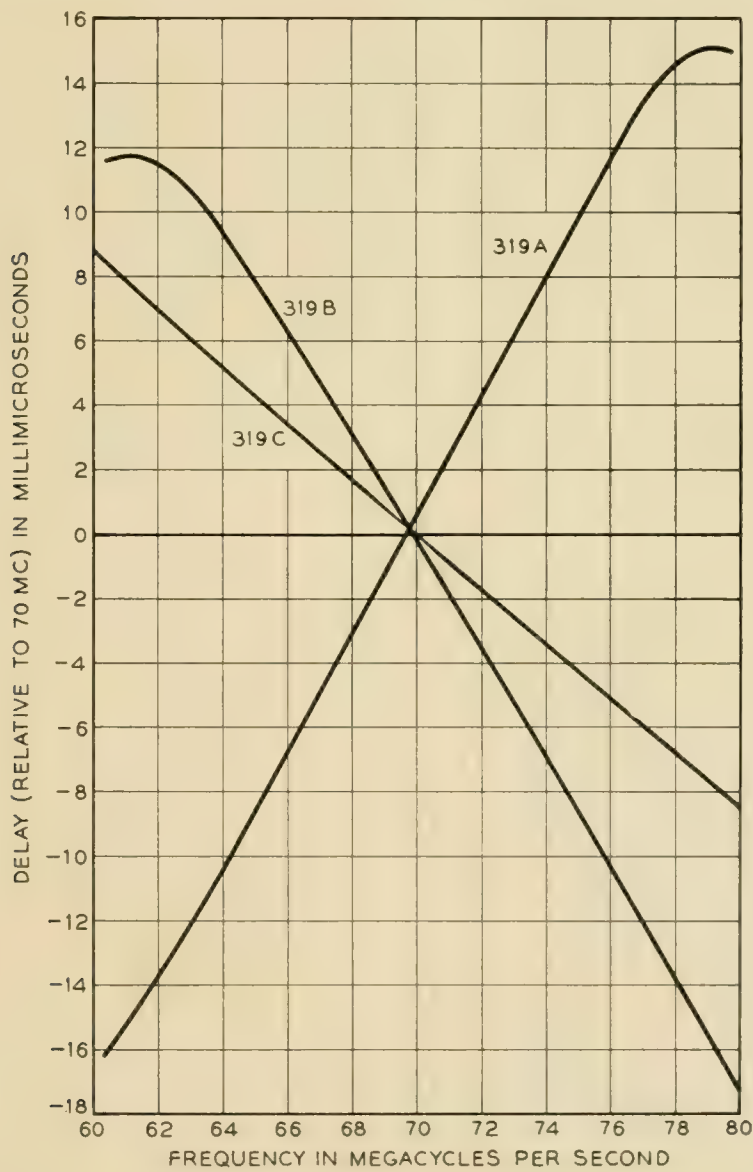


Fig. 2 — Delay characteristics of 319 type equalizers. Combinations of these are used at main stations to equalize delay slope of the system.

equipment produce a cumulative gain-frequency distortion which is noticeable in television circuits. Present practice is to correct for this by standard video equalizers after the FM signal has been demodulated to baseband. In connection with the experimental equalizing program to be described, parabolic gain equalizers operating on the FM signal before demodulation were used.

III. RESIDUAL DISTORTION

After correction of the known shapes discussed above, there remains a certain residual gain and delay distortion which results from a random summation of many minor sources. The shape of this distortion is not predictable, but its statistics are known. Examination of typical delay versus frequency characteristics have shown that these may be reasonably well approximated by six cosine terms: a 40-mc fundamental and the next five harmonics. Similar gain terms are needed. However, the gain and delay distortion, when examined within the 20-mc band of interest, do not have a minimum phase relationship. This is to be expected because of the presence in the system of the delay equalizers, which are non-minimum-phase networks, and of amplifiers with compression.

The magnitude of the residual distortion is small enough so that transcontinental TD-2 circuits provide television and telephone transmission of commercial quality. Some effects, such as cross modulation, are sufficiently marginal so that improvement would be desirable. To determine whether this could be achieved by improved gain and delay equalization, the development of an experimental adjustable equalizer was undertaken. The considerations outlined show that such an equalizer should approximate the desired characteristics with independent gain and delay terms of the harmonically related cosine type. Equalization to reduce cross modulation in telephone channels and differential phase in color television must be performed before demodulation of the FM signal to base band. The equalizer was, therefore, built to operate in the 60- to 80-mc IF band.

IV. TRANSVERSAL EQUALIZER

One method of obtaining independent control of the loss and delay characteristics of a network has been achieved in the transversal filter.² Equalizers have been designed on this principle for the equalization of television circuits.^{3, 4} This type of equalizer, referred to here as a transversal equalizer, provides a flexible means of synthesizing any loss char-

acteristic and any delay characteristic limited only by the number of harmonics that are provided and the range of each.

Basically, the transversal equalizer consists of a delay line with equally spaced taps, with a means for independently controlling the amount of signal fed through each of the taps to a summing circuit, as shown schematically on Fig. 3. The input signal is fed into one end of the delay line which is terminated at the other end. The center tap is fed to the output and forms the main transmission path.

The operation of the equalizer can best be described using the "time domain" analysis based on the theory of paired echos.⁵ Portions of the signal tapped off the "leading" or first half of the delay line will not be delayed as much as the main signal and will introduce leading "echos". Similarly, lagging echos can be obtained from the taps on the lagging or second half of the delay line. Combinations of both types of echos, either positive or negative as required, can be added to cancel out, to a first approximation, distortion present in the input signal.

This analysis can also be carried out in the frequency domain. To obtain a family of cosine loss versus frequency characteristics without any appreciable delay characteristic, equal leading and lagging echos of the same polarity are added to the main signal in the summing circuit. To obtain a corresponding family of cosine delay versus frequency characteristics without loss distortion, leading and lagging echos equal

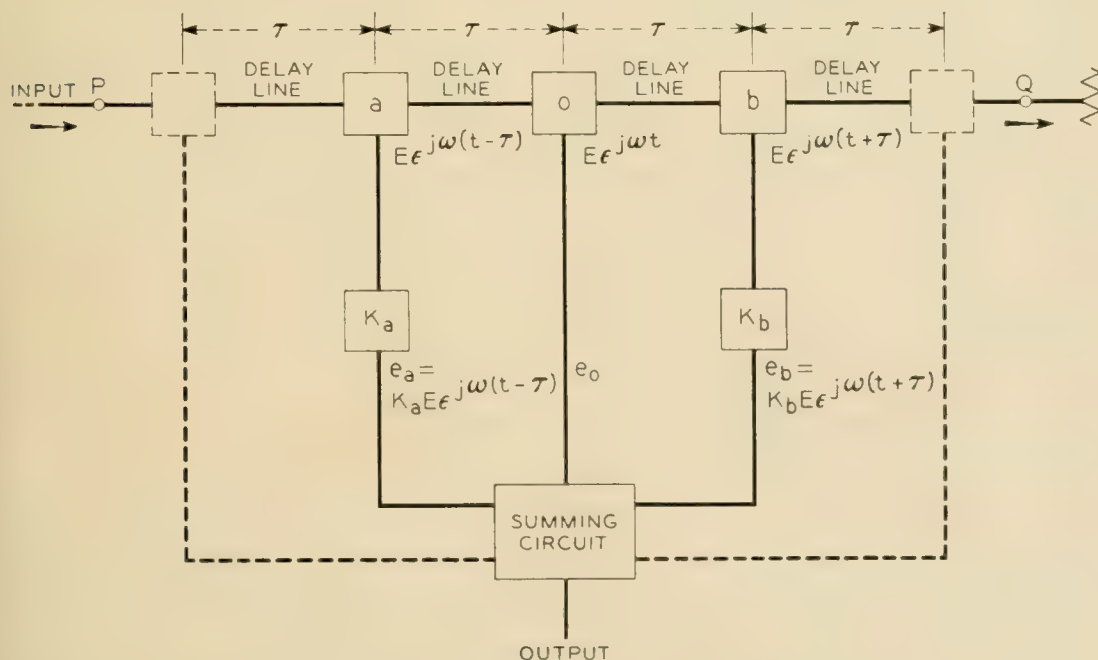


Fig. 3 — Block schematic of transversal equalizer.

in magnitude but of opposite polarity are added in the summing circuit to the main signal.

To achieve a practical equalizer for operation over the 60- to 80-mc band requires the following components: delay line or delay networks, means for tapping off small portions of the signal controlled both in amount and polarity, and a suitable summing circuit.

V. DERIVATION OF THE EQUALIZER CIRCUIT

A brief analysis of the operation of the equalizer will be given at this point as a basis for discussion of the method of tapping the signal and controlling the amplitude and polarity of the tapped portion.

Fig. 3 shows the basic delay line PQ as well as the means used for producing the main signal and a single pair of leading and lagging echos. The tap labeled "o" in the center of the line produces the main signal. The tap "a", being closer to the input, produces a signal which leads the main signal by time τ . The tap "b" produces a signal which lags the main signal by the same amount. The boxes " K_a " and " K_b " control the amplitude and polarity of the leading and lagging signals which are to be combined with the main signal to produce one term of the desired equalization characteristic. It will be shown that these three signals will provide one cosine gain term and one cosine delay term, both having the same period, but being independently controllable as to amplitude and polarity.

We will choose as our reference point for phase the main output signal, $e_0 = E\epsilon^{j\omega t}$. The output from tap "a" is then

$$E\epsilon^{j\omega(t-\tau)}.$$

After passing through box " K_a ", this becomes

$$e_a = K_a E\epsilon^{j\omega(t-\tau)}.$$

Similarly,

$$e_b = K_b E\epsilon^{j\omega(t+\tau)}.$$

Here the terms K_a and K_b are of the form

$$K = \pm \epsilon^{-\alpha}$$

where α is the attenuation in nepers of the box K . Note that $|K|$ is less than unity, assuming the box represents a passive network. Combining these two signals with the main signal, we have

$$e_r = e_0 + e_a + e_b = E\epsilon^{j\omega t}[1 + K_a\epsilon^{-j\omega\tau} + K_b\epsilon^{j\omega\tau}]. \quad (1)$$

Now it will be shown that, by the adjustment of the two parameters K_a and K_b , it is possible to realize independent control of a cosine gain term and a cosine delay term.

Since, in general, $K_a \neq K_b$, let us define

$$K_a = K_g + K_p$$

and

$$K_b = K_g - K_p. \quad (2)$$

Then

$$e_r = E e^{j\omega t} [1 + K_g (e^{-j\omega\tau} + e^{j\omega\tau}) + K_p (e^{-j\omega\tau} - e^{j\omega\tau})]. \quad (3)$$

Substituting the trigonometric form:

$$e_r = E e^{j\omega t} [1 + 2K_g \cos \omega\tau + j 2K_p \sin \omega\tau]. \quad (4)$$

Note that for K_a equal to K_b , the sine phase term is zero and that for K_a equal to $-K_b$ the cosine gain term is zero. Similarly, by proper proportioning of K_a and K_b , K_g and K_p may be assigned any desired values.

If we normalize (3) by setting $e_0 = E e^{j\omega t} = 1$, the expression in brackets can yield two vector diagrams which are useful in explaining the functioning of the equalizer. To obtain the diagram shown in Fig. 4(a), we have set $K_p = 0$. We then have a unit vector, representing the main signal, a leading echo $K_g e^{-j\omega\tau}$, and a lagging echo $K_g e^{j\omega\tau}$. The

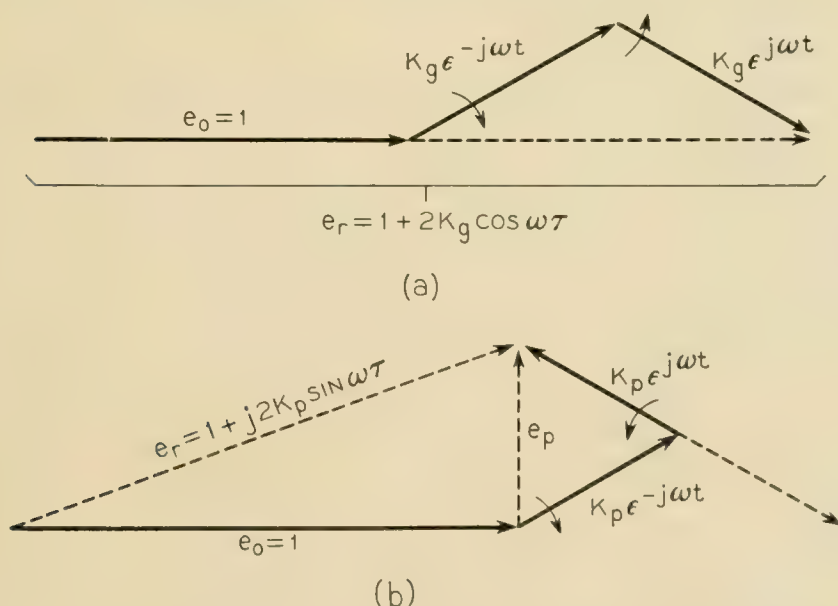


Fig. 4 — Vector diagrams of paired echos. (a) Equal echos of same polarity produce magnitude change without phase change. (b) Equal echos of opposite polarity produce a change in phase shift with a minor change in magnitude.

vector representing the leading echo rotates clockwise with respect to the main signal when the frequency increases, whereas the vector representing the lagging echo rotates counterclockwise by the same amount, and the resultant thus varies in magnitude but not in phase. The magnitude of the resultant is given, for this case, by the first two terms in parentheses in (4).

If, on the other hand, if we set $K_g = 0$, we have the three vectors shown in Fig. 4(b), identical with those in Fig. 4(a) except that the polarity of the lagging echo has been reversed. In this case, the two echos produce a resultant, e_p , which is in quadrature with the main signal. For small echos, e_r is thus shifted in phase from the main signal, with substantially no change in magnitude. The resultant in this case is given by the first and third terms in parentheses in (4). This gives a sinusoidal variation in the phase of the resultant. Since envelope delay is defined as $d\beta/d\omega$, where β is the phase shift through the circuit in question, the sinusoidal phase ripple will be seen to yield, after differentiation, a cosine delay ripple.

The period of the ripple can be seen from the above expressions to depend on τ , the delay between the leading and the main tap, and between the main tap and the lagging tap. Other pairs of echos, each pair symmetrically disposed about the main tap, but with different values for τ , will give transmission ripples of different periods. To provide a series of orthogonal terms, the values of τ must be integral multiples of a common value, normally that required to produce 180° phase shift across the band of interest.

A complete equalizer must, of course, sum up the various echos and the main signal, taking care that the delay between the tap and the summing point is the same for each echo and the main signal, that parasitic losses such as losses in cabling are the same for each path through the equalizer, and that any frequency characteristic in the tapping device or other parts of the equalizer is properly equalized out so that the over-all equalizer introduces a minimum of distortion of its own.

VI. DIRECTIONAL COUPLER

To reduce incidental distortion, it is desirable that the device used to tap the delay line for the main signal and the echos introduce substantially no discontinuity in the main line. The device chosen for this purpose is a directional coupler. It is shown symbolically in Fig. 5. The directional coupler is a four port device having properties similar to a

hybrid coil. Power entering one port divides (not necessarily equally) between two other ports, but none of it reaches the fourth, or conjugate port. In Fig. 5, the power entering at 1 divides between 2 and 4, that entering at 2 divides between 1 and 3, that entering at 3 divides between 2 and 4, and that entering at 4 divides between 1 and 3. Directional couplers inherently provide an impedance match at all four ports. Thus, such a coupler sets up no reflections in the main line. Its insertion loss in this line may be kept small by having nearly all the power entering at 1 come out at 2; then only a small fraction is diverted to 4. Coaxial directional couplers have been discussed in the literature^{6, 7, 8} and will not be dealt with in detail here.

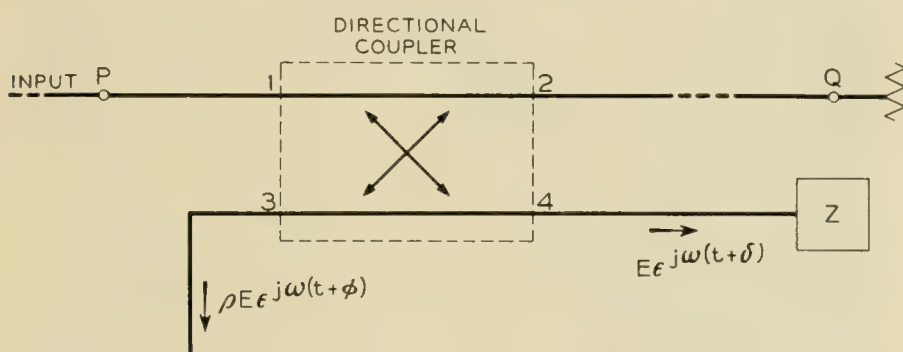


Fig. 5 — Diagram of directional coupler. Input signal divides between Ports 2 and 4 with no output at Port 3. Termination Z at Port 4 reflects some signal to Port 3, proportional to the reflection coefficient, ρ .

The coupler used here (J68333C) is one originally developed to measure reflections on IF transmission lines in the TD-2 system. The directivity of a coupler is defined as the coupling loss between main line and branch line in the undesired direction less the loss in the desired direction (loss from 1 to 3 less the loss from 1 to 4, for example). In the J68333C coupler, the directivity can be adjusted to exceed 45 db over the band of interest. This can be done by adjusting two screws, shown on model in Fig. 6, to obtain the optimum spacing between the coupling elements. The loss between the main line and the branch line in the desired direction is about 23 db at mid-band (70 mc), and decreases 6 db per octave with increasing frequency. The loss along one of the coupled lines (1 to 2 or 3 to 4) is very small.

Use has been made of the directional properties of the coupler in providing a simple means of controlling the amplitude and polarity of the tapped signal. Referring to Fig. 5, and keeping in mind the properties of the coupler, it will be noted that a small portion of the input signal appears at Port 4 of the coupler, but none at Port 3. If the impedance Z

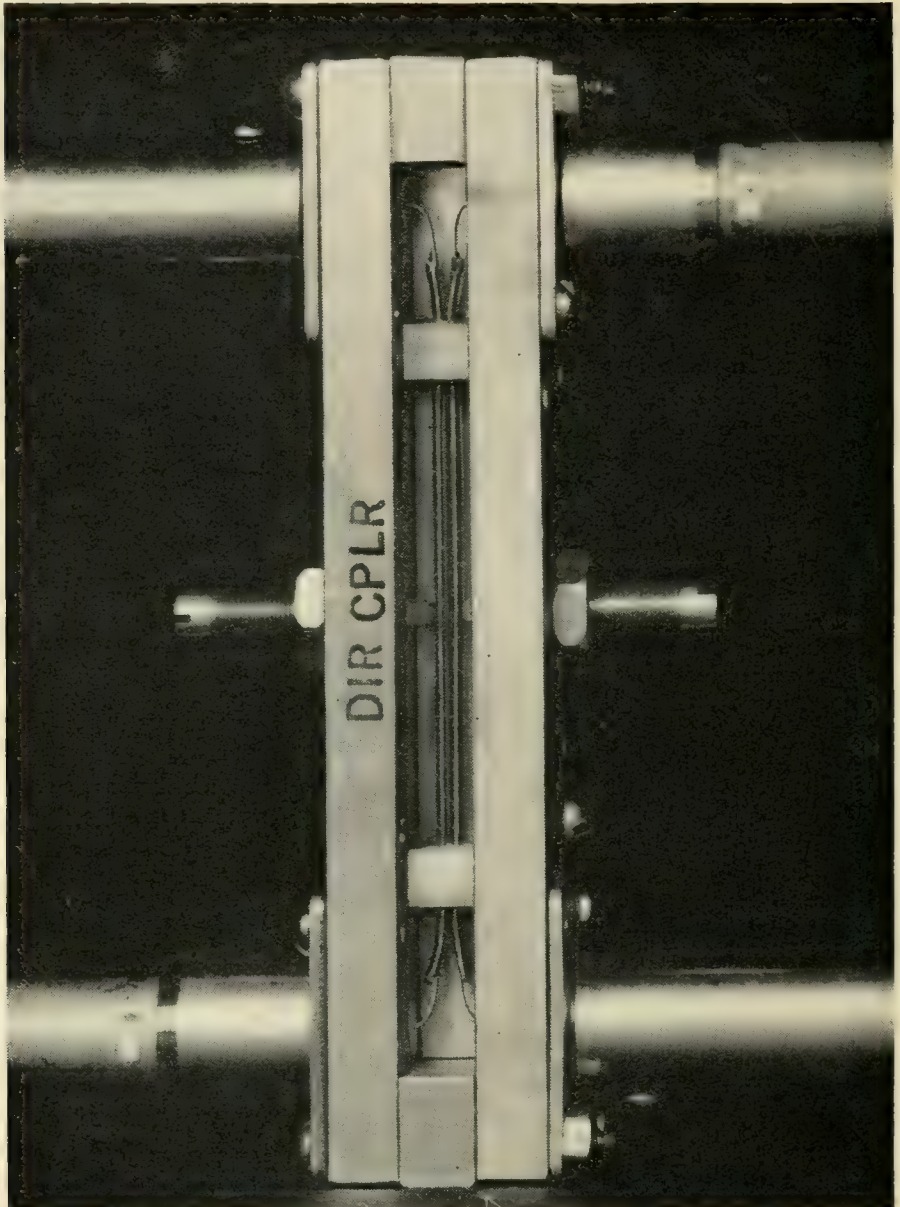


FIG. 6 — J-68333C directional coupler with cover plate removed to show coupling elements. Optimum spacing for maximum directivity can be obtained by adjusting screws.

matches the impedance seen looking into Port 4, all of the small signal will be absorbed in Z . If Z is an adjustable resistance, then a controllable portion of the small signal can be reflected back into Port 4, whence most of it will come out Port 3. A small portion of the reflected signal will emerge at Port 1, headed toward the input. This portion will be attenuated by twice the coupling loss between the main and the branch line plus the return loss of the reflection at Z , and can be made negligible. The interactions caused by the reflected signal on the main delay line entering previous couplers are also reduced by twice the coupler loss.

The coupler in Fig. 5 represents any one of the taps on the line PQ in Fig. 3. The signal emerging from Port 4 can be written as $E\varepsilon^{j\omega(t+\delta)}$, where δ is a time delay dependent on the location of the tap on the line PQ . The signal emerging from Port 3 is then $\rho E\varepsilon^{j\omega(t+\delta)}$, where ρ is the voltage reflection coefficient at Z , and is given by

$$\rho = \frac{R - R_0}{R + R_0}.$$

Here R is the value of resistance used to provide the impedance Z , and R_0 is the impedance seen looking into Port 4. An examination of the signal from Port 3 shows that it is the same as the signals e_a or e_b in Fig. 3, with the reflection coefficient ρ substituted for variable K_a or K_b in Fig. 3. Thus, it is seen that we may use the reflection at Z , variable by controlling the value of R , to perform the function of the box K in Fig. 3. Neglecting parasitic losses we may then write:

$$K = \rho = \frac{R - R_0}{R + R_0}$$

and

$$R = R_0 \frac{1 + K}{1 - K}. \quad (5)$$

This gives us the value of R to use for any desired value of K for any of the taps which derive echos, assuming the summing circuit has equal attenuation in all paths. In the case of the main central tap, the signal from Port 4 of the coupler is seen to be the same as the main signal e_0 in Fig. 3, and is used as such directly.

VII. METHOD OF ADJUSTMENT

The detailed design of a manually adjustable equalizer is materially influenced by the method to be used in the field for determining the setting of its controls. The present equalizer with 14 independent controls would present a complex problem of field adjustment unless special procedures were developed to simplify the adjustment. To adjust the equalizer, the radio circuit being equalized must be taken out of commercial service, so any reasonable measures to simplify the adjustment or reduce the time required are justified.

Two methods appeared to be feasible at the time the development was started. One would be to use existing gain and delay sweep test circuits. These present a visual display of the circuit gain or delay dis-

tortion versus frequency. These displays are not available simultaneously with present equipment. To adjust the equalizer controls using this equipment, it must be possible to adjust either gain or delay without affecting the other. Thus, all the gain terms can be adjusted in succession using the gain display. Then the procedure is repeated with the delay display, adjusting the delay terms. Since a combination of leading and lagging echos in equal amounts is required for this procedure, an arrangement of the controls to facilitate this is required. One way to achieve this is to use stepped rheostats with the steps proportioned to introduce equal amplitude changes in the echo voltage. With this arrangement, gain changes can be introduced by rotating the two switches corresponding to a pair of echos in the same direction an equal number of steps. Delay changes can be obtained by similar rotation in opposite directions. A further refinement consisting of mechanically ganging the controls is possible but this was not done on these experimental models.

The second method of adjustment would be to develop a special test set similar to the one developed for the L3 system.⁹ This could produce a meter reading proportional to the amount of gain and delay distortion present in the circuit. Successive controls could then be adjusted for minimum meter readings. Experience with the L3 system cosine equalizers has shown the desirability of continuously adjustable controls for such a method.

To test both methods under field-trial conditions, two versions of the equalizer were built — one with stepped rheostats and one with continuously adjustable rheostats.

VIII. COAXIAL RHEOSTAT

Since there were no available continuously adjustable rheostats satisfactory for operation at 70 mc, a special rheostat was developed for

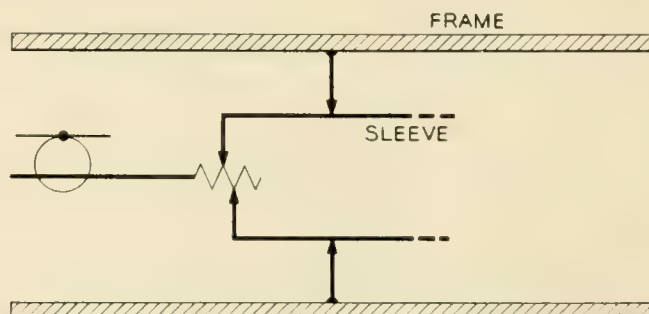


Fig. 7 — Schematic of coaxial rheostat. Moveable sleeve changes position of inner contacts touching ceramic rod, changing resistance. Fixed outer contacts maintain constant path length to frame.

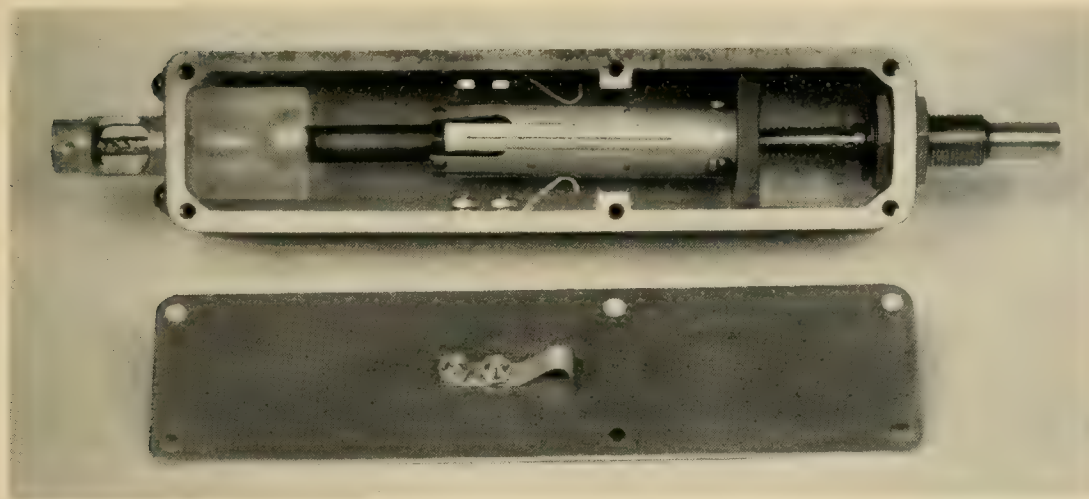


Fig. 8 — Model of coaxial rheostat with cover removed.

this purpose. It employs a ceramic rod, $\frac{1}{4}$ inch in diameter, coated with a pyrolytic carbon film, as a center member of a coaxial structure. A metal sleeve which is moved longitudinally by a lead screw carries sliding contacts along the rod. These parts are supported inside a rectangular housing which forms the outer conductor. A second set of fixed contacts attached to the rectangular housing makes contact with the sleeve. This arrangement maintains a substantially constant length of path from the input end of the rod to the housing, which forms the ground, independent of the position of the sleeve. The schematic of the rheostat is shown in Fig. 7. A model of the rheostat is shown in Fig. 8.

To obtain uniform adjustment of amplitude in decibels, a resistance that varies exponentially with length or with rotation of the lead screw is required. Such a resistance characteristic is realized by varying the thickness of the carbon film along the rod. This produced a total resistance which varied from 20 ohms at the low setting to about 350 ohms at the high resistance setting. After an initial wearing-in period of 1,000 cycles of moving the contacts over their full travel, the resistance was changed less than 1 per cent by another 9,000 cycles. This amount of wear is estimated to be greater than that encountered in twenty years of normal operation.

The housing and rod were dimensioned to form a 75-ohm transmission line. Measurements of the impedance at the input connector, made at frequencies from 60 to 80 mc, showed that this impedance can be approximated by a resistor terminating 6.7 cm of 75-ohm coaxial cable. For the 75-ohm setting, the reflection coefficient of the rheostat is less than 2 per cent across this frequency band.

One model of the equalizer was completely equipped with these rheostats. By allowing for the equivalent length of cable within the rheostat, an essentially pure resistive termination was obtained.

IX. OTHER COMPONENTS

Other components required for the equalizer included stepped rheostats, delay line or delay networks, a suitable summing network, and a loss equalizer.

The stepped-switch rheostats were made from standard switch parts with eleven positions. Deposited carbon resistors, 205D, were used for the steps. The mid-position corresponded to the circuit impedance level, 75 ohms. The other steps were arranged to provide equal increments of echo amplitude measured in decibels. With careful control of lead lengths, special shielding and a coaxial cable connector, satisfactory control of the return loss of this rheostat was obtained.

Resistance pads were added to the switch assemblies associated with each of the echo terms. The loss of each pad was determined so that the corresponding term would have the desired maximum amplitude. In addition, the losses of the pads associated with the leading echo terms were increased to compensate for the midband loss of the delay line between the leading coupler and the corresponding lagging coupler. This insured that the two echos would have equal amplitudes.

The delay required between taps in the delay line is 0.025 microsecond, corresponding to a change in phase shift of 180° from 60 to 80 mc. In order for the equalizer cosine characteristics to have maxima at the band edges, the total phase shift at 60 and 80 mc must be successive integral multiples of 180° . Since the phase shift of coaxial patch cable is closely linear and proportional to length, it could be used for the delay line. Lumped-element delay networks consisting of two or more all-pass sections are a feasible alternative and would reduce the over-all size and weight. In view of the additional development effort involved to produce these and the experimental nature of this equalizer, it was decided to use coaxial patch cord. The type selected, 728A cable, has a polyethylene dielectric and is tested during production for return loss in the 50- to 95-mc band. The length required for each section is about 15.6 feet. This much cable has a loss of about 0.3 db at 70 mc.

It was originally proposed to use a series of directional couplers for summing the echo voltages with the main signal. This would provide additional isolation between terms. However, tests on a preliminary model indicated this isolation was not required in this application. In-

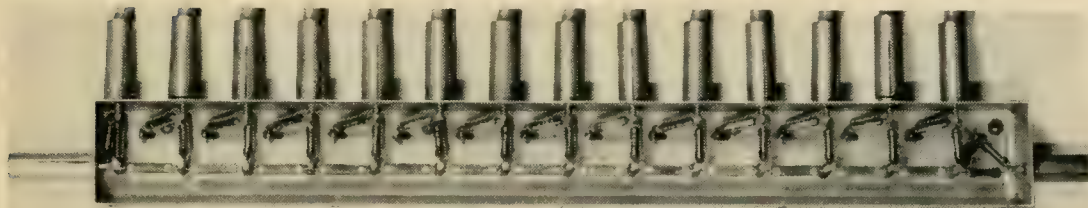


Fig. 9. — Summing network with cover removed. Main signal input is at right, echo signal inputs on top, and output is at left..

stead, a resistance summing network was developed using deposited carbon resistors. An L-pad is used in each echo path and a series resistor is added to the main path to preserve the 75-ohm impedance level, introducing a main path loss of about 0.4 db per tap. A model with cover removed is shown on Fig. 9. The main signal is introduced at one end of the structure, the echo voltages are connected along the side and the sum is taken off the other end. The return loss measured at any of the connectors with the others terminated was of the order of 40 db over the 60- to 80-mc band.

An attenuation equalizer is required to make the transmission through the main path constant. This path consists of about 108 feet of 728 cable, the straight-through loss of six couplers, and the coupling loss of the main coupler. The net distortion over the band is a slope of about 1.5 db and is corrected for by a constant resistance equalizer. Return losses exceeding 34 db were obtained over the frequency band.

X. ASSEMBLY

All the components were mounted on the rear of a standard relay rack panel. The rheostat controls are arranged on the front of the panel as shown on Fig. 10. This is a front view of the completed equalizer. The controls for the leading echo terms are on the left and for the lagging echo terms on the right. They are arranged vertically in numerical order with the first terms (shortest time separation from the main signal) at the top.

The rear of the panel is shown on Fig. 11. The directional couplers are mounted horizontally in two vertical columns. The cables forming the delay line sections are terminated in a plug and a jack and these are inserted in successive couplers, from the second port of one to the first port of the next. The third port of each coupler is connected through a short cable to its corresponding rheostat assembly. The fourth port of each coupler is connected to the summing network. An exception is

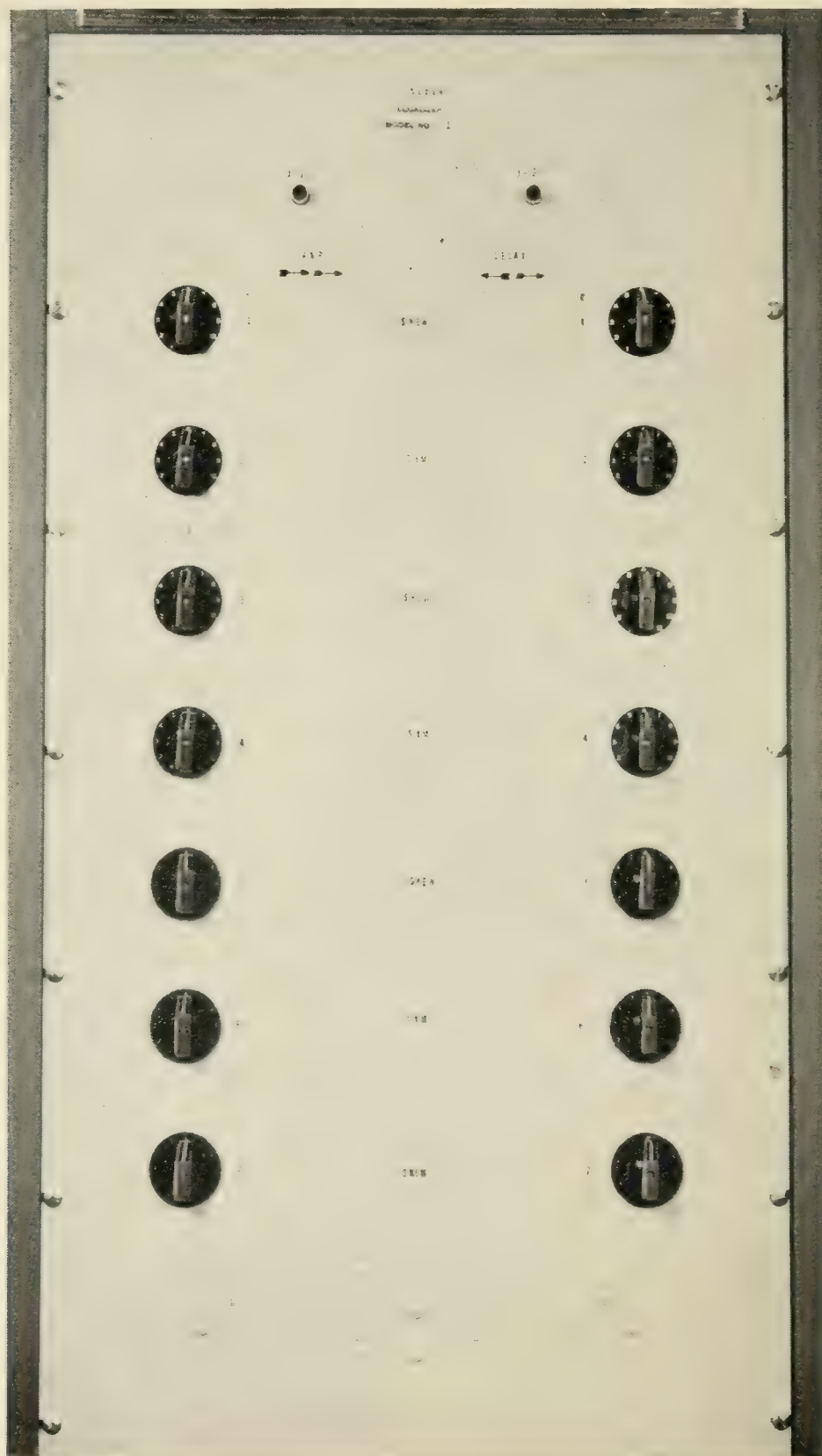


Fig. 10 — Front view of equalizer, showing rheostat controls. Leading echo controls are on left, lagging ones on right. First harmonic controls are at top.

the middle coupler, the fourth port of which is terminated in 75 ohms and the third port connected to the summing network.

The envelope delay in the cables connecting each coupler to the rheostat and to the summing network appears as delay for the particular echo path. Since these delays are not negligible compared to the 25 millimicrosecond delay between echos, the cable lengths were controlled so that the same amount of additional delay was introduced into each path including the main path.

XI. ADJUSTMENT AND PERFORMANCE

After the equalizer was assembled, the length of each of the cables connecting the couplers to the summing network was adjusted so that

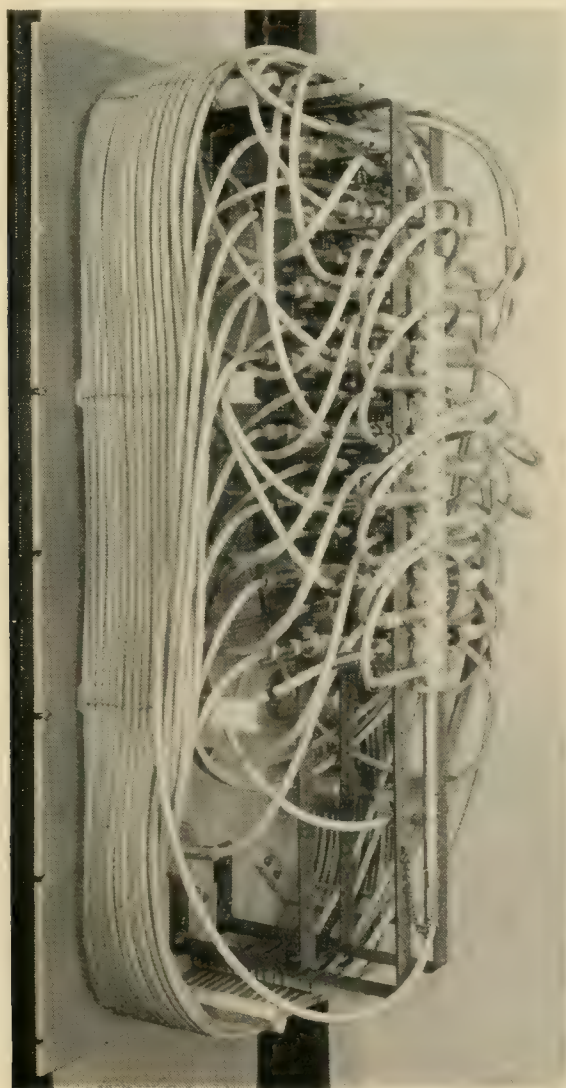


Fig. 11 — Equalizer with cover removed. Cables wound outside frame are the delay line sections. Summing network and directional couplers are in center. Rheostat cases are mounted on panel with coaxial connection in rear.

the zeros of the cosine shape occurred at the proper frequencies, as observed when the associated rheostat was set at maximum and minimum positions, with all other rheostats set at midrange, or no-echo, setting.

Some reflections were present in the main signal path as evidenced by ripples in the gain characteristic when all rheostats were set at midrange, corresponding to the "flat" loss condition. These reflections were reduced to some extent by minor readjustments of the balancing screws on the directional couplers. The over-all flat gain characteristic obtained after these adjustments is shown on Fig. 12.

This figure also shows the seven gain characteristics obtained when each pair of rheostats is set for maximum gain. The markers on the reference trace correspond to the band edges. A sharp gain bump resulting

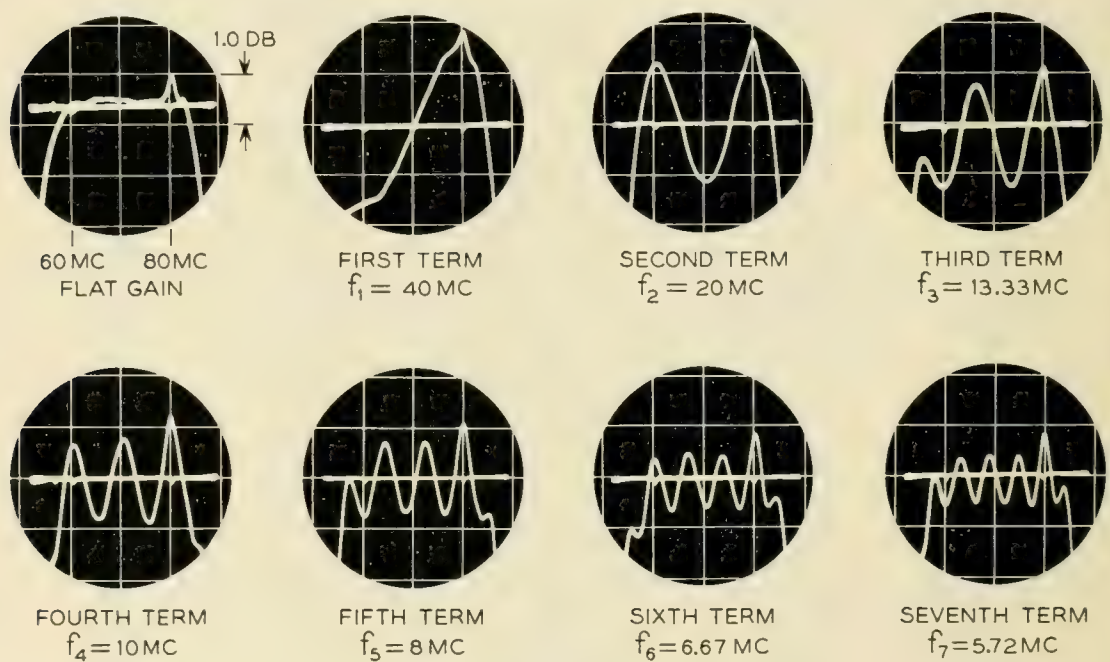


Fig. 12 — Measured gain characteristics of equalizer.

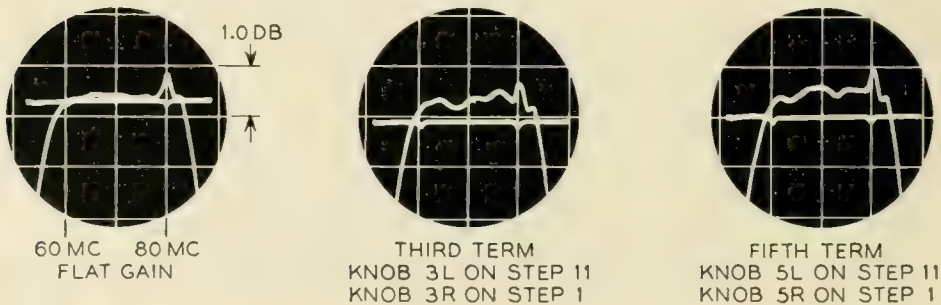


Fig. 13 — Gain change introduced by changing delay terms from zero to maximum. Left, normal case. Middle, third harmonic term at maximum. Right, fifth harmonic term at maximum.

from reflections on the delay line occurs just above 80 mc and distorts each characteristic near this frequency. The delay characteristics obtained closely resemble the corresponding gain characteristics. The delay characteristic with all rheostats on midrange was flat to within about ± 1.5 millimicroseconds.

Another measure of the performance is the amount of interaction between gain and delay characteristics. As shown on Fig. 13 the gain changes less than 0.2 db when the third or fifth harmonic delay terms are set at their maximum values of 11 and 12 millimicroseconds, respectively. Similarly, when a pair of rheostats are set for the maximum gain characteristic, the effect on delay is of the order of two millimicroseconds. These results, which are typical, indicate the interaction effect is of the order of 20 per cent using one neper of gain distortion ripple as equivalent to a ripple of one radian of phase shift amplitude. The effect of this interaction on the field use of the equalizer is to require a second round of equalization to correct for interactions after the gross distortions in a circuit have been equalized.

XII. FIELD EXPERIMENTS

Models of this equalizer were installed in two channels of the TD-2 system between Denver, Colorado, and Omaha, Nebraska, early in 1956. This route is about 500 miles long and includes 18 microwave links. One equalizer was installed at the center of the route and a second one at the receiving end. The results of a typical set of characteristics obtained are shown in Figs. 14 and 15. The first shows the delay characteristic of the whole channel measured at the receiving intermediate

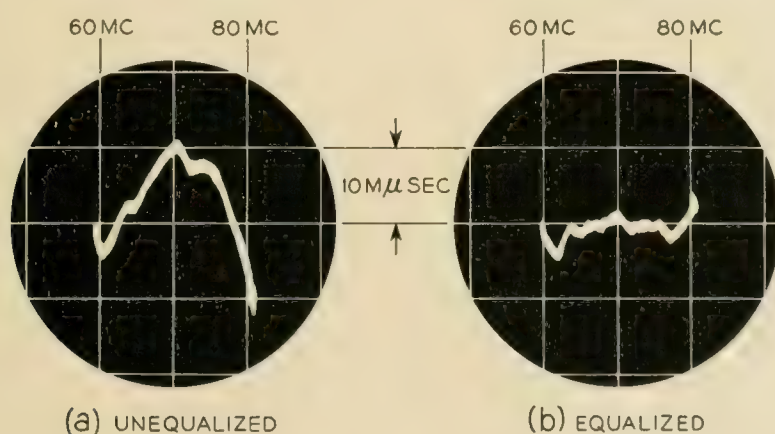


Fig. 14 — Measurements on Denver-Omaha route. Envelope delay distortion measured at intermediate frequency point. Left, unequalized; right with equalizer adjusted.

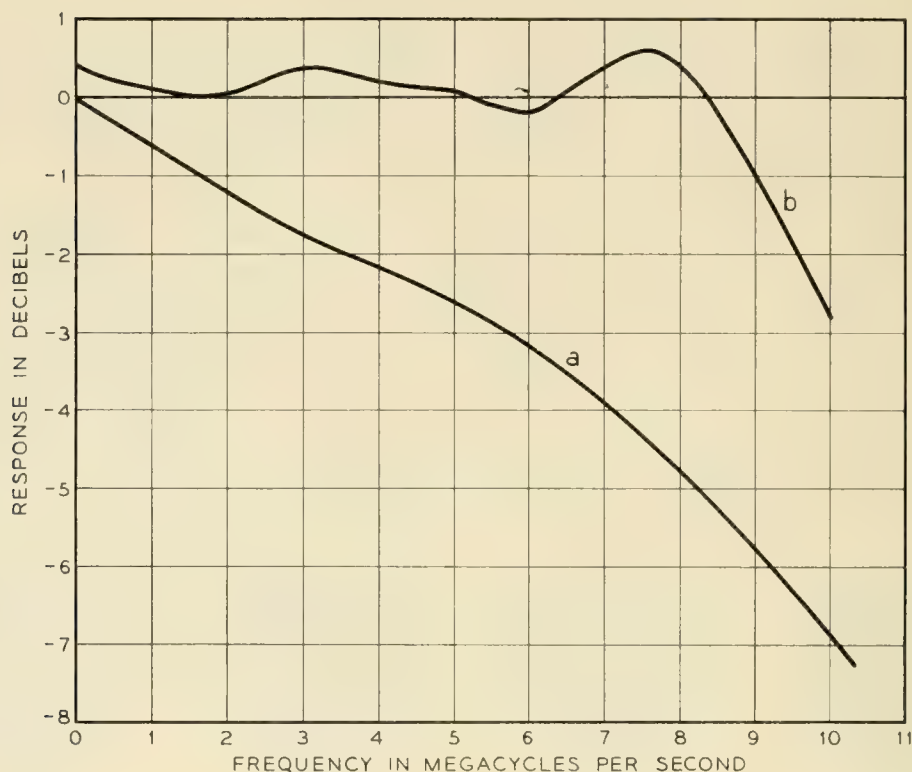


Fig. 15 — Measurement on Denver-Omaha route. Transmission characteristic measured at video (a) unequalized, (b) with equalizers adjusted.

frequency point. The “unequalized” characteristic is the circuit delay distortion immediately after standard line-up procedures without video equalization. The “equalized” characteristic shows the same circuit distortion corrected by the addition of the transversal equalizer. The first two delay terms of this equalizer were supplemented by the use of 319 type (linear slope) and 315A (parabolic) equalizers. The second loss term was supplemented by the use of experimental fixed parabolic loss equalizers. An attempt was made to obtain the best performance over the center 10 mc of the band, corresponding to the first order modulation band.

The gain distortion was adjusted on a demodulated video basis as this was the only type of sweep gain circuit available. As shown in Fig. 15, the unequalized circuit had more than 0.5 db per mc of slope up to 10 mc. The addition of the equalizer produced a flat band to about 8.5 mc.

The effect of this equalization on cross modulation, measured by simulating the message load with a flat band of noise, is shown in Fig. 16, for two sample channels. In these curves, normal drive represents the noise load whose power is 12 db below the power in a sine wave giving 4-mc peak deviation of the TD-2 carrier. Channel A showed 5-db im-

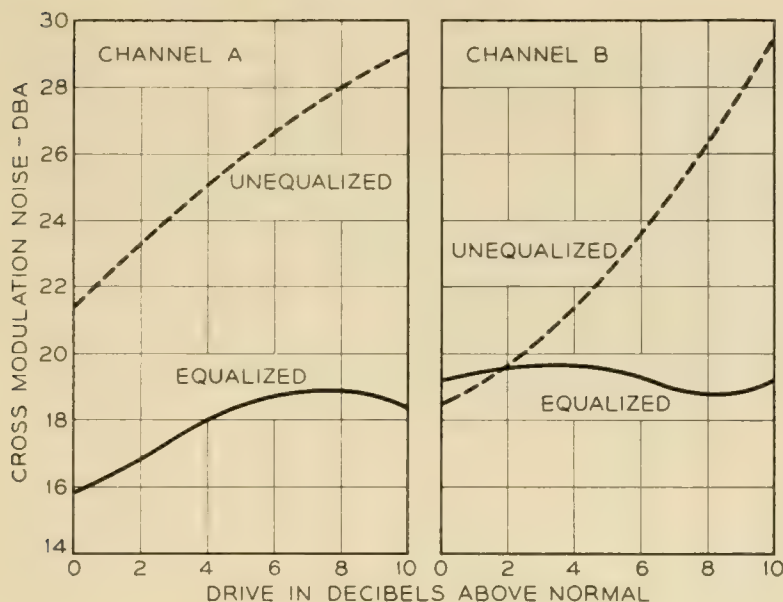


Fig. 16 — Effect of improved equalization on cross modulation noise, Denver-Omaha route. Measurements on two channels, referred to -9 db transmission level.

provement at normal drive, and a somewhat greater improvement at higher drives. Channel B, however, showed a slight degradation at normal drive, which becomes a 9- or 10-db improvement at high drive. In the case of the latter channel, the improvement at normal drive was limited by the presence of a delay ripple, due to waveguide echoes, which was of too short a period to be equalized by the present equalizer.

The realization of such improvements in a working system is limited by such waveguide echoes, which are not stable enough for ready equalization, as well as by other instabilities in the transmission characteristic which have been attributed to antenna and air path effects.

ACKNOWLEDGMENTS

Several of the authors' colleagues contributed materially to this development. C. H. Dagnall developed the 315 and 319 type equalizers. A. H. Volz and J. K. Werner developed the coaxial rheostat, based on suggestions by S. Bobis. Mrs. Grace L. Ebbe was responsible for much of the equalizer design and the assembly and testing of models. J. G. Chaffee and I. Welber directed the field tests.

REFERENCES

1. A. A. Roetken, K. D. Smith, and R. W. Friis, The TD-2 System, B.S.T.J., **30**, pp. 1041-1077, Oct., 1951.

2. H. E. Kallman, Transversal Filters, Proc. I.R.E., **28**, pp. 302-310, 1940.
3. J. M. Linke, A Variable Time Equalizer for Video Frequency Waveform Correction, Proc. I.E.E., **99**, Part IIIa, pp. 427-435.
4. R. S. Graham and R. V. Sperry, Equalizer, U. S. Patent No. 2,760,164.
5. H. A. Wheeler, The Interpretation of Amplitude and Phase Distortion in Terms in Terms of Paired Echos, Proc. I.R.E., **27**, pp. 359-385, 1939.
6. G. D. Monteath, Coupled Transmission Lines as Symmetrical Directional Couplers, Proc. I.E.E., Part B, No. 3, **102**, pp. 383-392, May, 1955.
7. B. M. Oliver, Directional Electromagnetic Couplers, Proc. I.R.E., **42**, p. 1686, Nov., 1954.
8. B. C. Bellows, Directional Coupler, U. S. Patent No. 2,679,632.
9. R. W. Ketchledge and T. R. Finch, The L-3 System - Equalization and Regulation, B.S.T.J. **32**, pp. 833-878, July 1953.

Transmission Aspects of Data Transmission Service Using Private Line Voice Telephone Channels

By P. MERTZ and D. MITCHELL

(Manuscript received February 2, 1957)

An exploration is reported of the possibilities of a moderately fast data transmission system to use private line message facilities. A comparatively conventional system was desired to permit expeditious application. An auxiliary "word start" signal was necessary for the system considered.

Transmission characteristics of a number of arrangements were examined. These included several exploratory AM vestigial sideband systems (using a spectrum similar to telephotography), double sideband AM systems, various telegraph and other multichannel systems.

It was concluded that the 1600 bits per second usable over the AM vestigial sideband arrangement was about as fast as could be expected with the data system contemplated. This transmission is not as "rugged," with respect to impulse noise and sudden level changes, as other slower arrangements, but it is expected to be satisfactory. It will require delay correction, and simple methods are considered for carrying this out.

I. INTRODUCTION

The Bell System has been approached on a number of occasions in regard to the transmission of computing machine and similar data over its telephone circuits. This has reached the point where specific possibilities for private line data transmission have been given serious consideration.

The telephone network was developed for speech transmission, and its characteristics were designed to fit that objective. Hence, it is recognized that the use of it for a distinctly different purpose, such as data transmission, may impose compromises both in the medium and in the special service contemplated.

A short time ago the authors were assigned the problem of examining possibilities for such an adaptation aimed at high speed, and exploring

the nature of the transmission compromises that might be needed. As a result, a variety of data transmission systems in different stages of development have been investigated including some telegraph systems. Certain conclusions have been reached regarding their suitability for use over private line telephone facilities. (By "private line" facilities are meant facilities leased to the subscriber on a more or less permanent private basis, and that are not set up by operators at a telephone central office switchboard.)

These conclusions are summarized below. In the later text there is included the background material for the conclusions which includes a brief characterization of the facilities. Finally there is included a discussion of the line treatment which may be needed for the best transmission.

It should be emphasized that not all possible data systems for use over telephone circuits are covered. The problem considered covers particularly some recently proposed applications, for which the need of a relatively high bit rate is important. Also, only the more promising comparatively conventional systems, which have been relatively well tested and can be readily applied, are considered. More radical designs are conceivable but they would require more extensive investigation before conclusions could be reached concerning them. It is clear also that the designs involved in the choice of a system are determined by the type of service it is to provide.

1.1 *Conclusions*

It is concluded that about the fastest transmission of data which can be accomplished with the present art over message-type telephone facilities is obtainable with an amplitude modulated vestigial sideband system. Such a system will be presented in some greater detail below. Its frequency spectrum is similar to that of a telephotograph signal² and a number of the transmission problems involved are the same.

This system will provide about 1600 binary digits of data (or "bits") per second, but it requires some special selection and considerable treatment of many types of circuits. This treatment is necessary to reduce noise, particularly of the impulsive type, and to correct for delay distortion.

Such a system is therefore considered suitable where a high bit rate is essential. It will be developed in the text that beyond the matter of the delay correction and other treatment of circuits required, the vestigial sideband process imposes a certain signal-to-noise penalty. A further

signal-to-noise penalty comes from the method of multiplexing an auxiliary channel needed for word start indications.

Where a lower bit rate is acceptable, a 750 bit per second amplitude-modulated, double-sideband system which is now available, and has been tested extensively, is somewhat more rugged. It permits operation over the great majority of telephone facilities and will also be described in more detail.

The most rugged system considered uses frequency-shift transmission. This form of transmission has been used extensively for some time on a multiple-channel, voice-frequency telegraph basis in which each channel is capable of 46 to 74 bits per second. When used in this form to handle high speed data signals, it requires relatively complicated terminal devices because the several channels have to be merged to provide one high speed data system. However, when frequency shift is used as a single channel over an untreated telephone message facility the system promises to be relatively simple and to give a total possibility up to 1,200 bits per second according to the type of facility.

1.2 *Summary Table*

These findings are concisely grouped in a summary table, Table I. These entries are based upon present knowledge and are believed to be reasonably accurate, although estimates regarding impulse noise need more extensive checking, particularly in the case of the broad-band, frequency-shift system.

The table compares relative estimated performances of three broad-band systems among themselves, and with a subdivided channel (or telegraph) system. The performances considered cover the effects of noise and delay distortion, and the bit rates considered achievable in a 1,000-cycle band and in a 300- to 2,800-cycle telephone band. Some crude approximations covering speech are also given as a matter of interest.

The noise performance is given in terms of relative total power capacity required in the line for a given error rate in the presence of a given noise. The double sideband system is taken as reference. Allowance is made in the multiple channel or telegraph system for occasional peaking caused by temporary unfavorable phasing. In this part of the comparison a 12 channel system has been assumed. This is about as many channels as can be used on a telephone facility that has heavy impulsive noise.

The delay distortion figures represent some present ideas on good engineering design in the allowable impairment of signal-to-noise ratio.

TABLE I—SUMMARY OF VARIOUS DATA SYSTEM CHARACTERISTICS

Use in Band	Relative Peak Power Required in Line for Equal Performance* in Presence of Noise, db			Approx. Max. Delay Distortion in Millisec. Allowable in Band (for 3 db noise penalty)		Approx. Bits per Second in 1000 Cycle Band (25-db S/N Random Noise for DSB)	Approx. Max. Bits per Second in Telephone Facility 300-2,800 Cycles
	SF	Random	Impulse	Within One Channel	Over 1000 Cycle Band		
Broadband VSB	+6	+6	+6	± 0.4	± 0.4	1000	1600**
Broadband DSB	0	0	0	± 0.55	± 0.55	700	800 to 1400**
Broadband FS	-3	-3	-7†	± 0.5	± 0.5	650	750 to 1200**

Telegraph Comparison

43A1 Channel	+10†† (—————)	0†† 12	-10†† Channels	± 5 (—————)	± 60 (—————)	450 (6 Channels)	1100 (15 Channels)
--------------	------------------	-----------	-------------------	--------------------	---------------------	---------------------	-----------------------

Speech Comparison

Speech	+10	+10	+10	± 10			50
--------	-----	-----	-----	----------	--	--	----

* High grade performance; i.e., less than 1 error per 100,000 bits.

** High figure assumes accurate delay correction and control of nonlinearity.

† Depends on precise line-up of filter and carrier.

†† Allowance for peak factor of 12 frequencies.

VSB = Vestigial sideband; DSB = Double sideband; FS = frequency shift.

This impairment is here assumed to be about 3 db. A rather larger impairment has on occasion been assumed by other authors.³ In the case of the telegraph system some arrangements for merging the signals in parallel channels into a single high speed channel require a certain exactness in timing correlation. The need for this may be overcome by the use of a small amount of data storage in the receiver of each channel. The delay distortion figures quoted in the table assume no need for this timing exactness. As in the other part of the comparison twelve channels are assumed.

The last two columns indicate the bit rate that can be expected of the various systems, first per 1,000 cycles of band, and second for a telephone facility of somewhat narrow (but frequently encountered) bandwidth. Some of these figures assume a careful control of delay distortion and of nonlinear distortion. In the 1,000-cycle band only six channels of the telegraph assumed may be accommodated. In a 2,500-cycle band the number can be extended to 15. The band that can actually be used for telegraph, over a given facility, depends upon the nature of that facility.

Some comparison data are indicated for telephone speech. The bit rate given assumes particularly message communication and not finer shades of artistic expression. For this a collection of phonemes (or elementary sounds) in the fifties, needing six bits for identification, and a typical speech rate of eight phonemes per second, require 48 bits per second. A few bits are also needed for pitch indication, and the figure is rounded to 50.

The low bit rate obtained for telephone speech communication indicates considerable redundancy in the speech signal sent over the telephone channel. This suggests why substantial transmission impairments can be tolerated without destroying the intelligibility of speech, as compared with telegraph or data signals.

1.3 *Nature of Study*

The procedure followed has consisted of examining systems of binary or similar signal transmission which appeared suitable for the sending of data. The systems studied are listed here, and some description of them is given later in the text. As noted, part of the information has come from outside the Bell System.

1. Exploratory 1,650 bit per second vestigial sideband system studied at Bell Telephone Laboratories.

2. Exploratory 1,600 bit per second vestigial sideband system studied at Lincoln Laboratory of M.I.T.¹

3. Exploratory 750 bit per second double sideband system, of general type reported by Horton and Vaughan.³

4. Voice frequency telegraph channels.⁵

5. "Polytonic" signaling system reported by Lovell, McGuigan and Murphy.⁴

The transmission problem of applying these various systems to the various types of telephone message facilities employed for private line service has been considered. These are also listed, and a brief characterization of them is given in the text.

1. Voice frequency circuits, over cable and open wire.

2. Type-C carrier circuits for open wire.¹⁶

3. Type-N carrier circuits for cable.⁶

4. Broad-band carrier systems¹⁵ using A channel banks for paired cable, coaxial cable, open wire, and microwave radio.

5. Other broad-band carrier systems.

1.4 *Outline of Paper*

As a result of the study, some recommendations were made. The discussion of this material is presented in Section II. It covers first the sort of service, with comparatively high bit rate, for which there has been user demand, and which can be furnished over private line facilities. Secondly, the recommendations cover the broad features of the signal characteristics which appear promising for use over such facilities. These recommendations are at the present time evolving into an exploratory system. The recommendations themselves, together with some general remarks, are presented in Sections 2.0 and 2.1. The background material on which these were based, and which covers consideration of the five systems which have been listed, is presented in Sections 2.2 to 2.5.

The discussion on the nature of the problems involved in transmission of the signals over the telephone plant to be used (which above has been listed in five categories) is finally covered in Section III.

II. SYSTEMS OF BINARY SIGNAL TRANSMISSION

2.0 *General Remarks*

There are a number of arrangements in the present art, both experimental and commercial, which are capable of sending digital data information under something like the conditions required for the service considered. Study of these has led to a set of recommendations which are outlined below. The specific arrangements are then discussed in more detail. The description covers only the essential transmission characteristics of the arrangements.

The arrangements generally divide into two groups, those which are essentially short-pulse single channel (though some of these may include a slow auxiliary channel), and those which use frequency division multiplex channels and therefore employ longer individual pulses. One important advantage of the multiple channel systems is noted in Section 3.3 as consisting of an increased immunity against impulse noise.

The single channel group shows much similarity among the systems. The principal difference is that of bit rate. The faster systems use vestigial sideband transmission, at the expense of a certain increase in vulnerability to noise as compared with those which use double sideband transmission, and also at the expense of a more general need for delay distortion correction because of the speed.

The systems in this group which use vestigial sideband are expected to be able to operate, with a very low error rate (some 1 error per 100,000

bits), over telephotograph type circuits, which employ delay distortion correction. This is a "normal" condition error rate, and does not include abnormal circuit conditions, such as static, trouble conditions, etc. It also assumes a more or less even error distribution.

One of these systems gives very desirable word start indications by using an additional level of signal. "Words" are groups of signal elements of fixed total number each. A distinctive separation signal greatly simplifies their recognition at the receiver, particularly after short line interruptions.

As the price, however, of both the speed and the third level signal indication, this system is more vulnerable than the others to impulse noise, of the kind encountered in hitherto installed N1 carrier and open-wire circuits, which have not been treated for impulse noise. It is also more vulnerable than the other systems to sudden level changes in the received signal.

The multiple channel group is generally characterized by a lower bit rate. Of all of them only one, the frequency-shift carrier telegraph, shows capability of consistent operation over untreated N1 carrier. On this type of carrier the frequency-shift telegraph permits a bit rate of the order of only half that obtainable with the vestigial-sideband systems. This telegraph system performs much better over other types of telephone facilities but even then its bit rate is only about three fourths that desired.

2.1 Digital Data Transmission Service

2.1.1 Some Desirable Major Requirements

1. Transmission of a maximum of 1,600 bits per second with an error rate not to exceed one in every 100,000 bits (or once per minute).

2. Applicability to most telephone facilities for private line use. Some selection of circuits and some treatment of those selected will be acceptable. The circuits carrying data are to be one-way terminal circuits only (i.e., not to be linked in tandem). It is expected that for the bulk of the service the circuits would not be likely to run over some 200 to 300 miles in length, but a small number of 3,000-mile circuits is considered possible.

2.1.2 Characteristics of the System

The system which has been considered and recommended as promising for the service outlined above has the following essential characteristics:

1. Carrier at 2,000 cycles, with vestigial band extending up to some 2,400 cycles. Lower nominal effective band (half the bit rate in width)

extends down to 1,200 cycles, and roll-off band down to about 1,000 cycles. The frequency space above the vestigial band and below the roll-off band is essentially "dead" space, free of signal energy.

2. The signal comprises three components:

- (a) *synchronization*, (or start), to indicate word separations,
- (b) *data*, or the actual information, and
- (c) *timing*, to mark out the successive bit intervals in the data.

The signal is represented in Fig. 1 as the envelope of a carrier. The synchronizing signal has an amplitude of one and a duration of one signal element. Data spacing signals have an amplitude of about 0.63 and a duration each of one signal element. Data marking signals have an amplitude of about 0.25 and also one signal element duration. It will be noted that this is an "upset" signal, in which spaces have more power than marks. The two signal elements immediately before, and again the two signal elements immediately after, the synchronization signal are always to be spacing. A timing signal is not actually sent, but instead the synchronizing signal is used to pull the phase of a highly

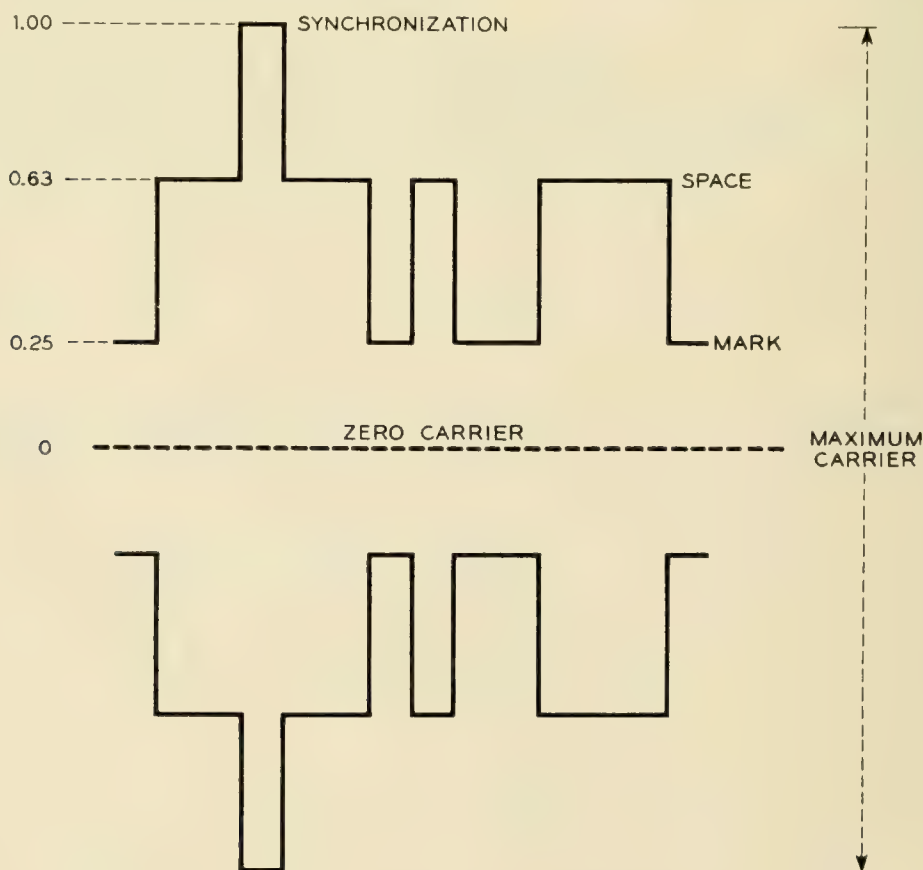


Fig. 1 — Multiple level signal.

accurate oscillator or "clock" at the receiver into the proper phase relation at the beginning of each word.

The phase of the oscillator is readjusted slightly as necessary at the beginning of each word but this adjustment is purposely made sluggish so that an occasional noise burst will not throw timing out badly. Thus the timing information is, in effect, by its repetitive characteristic, highly redundant. The receiving device is capable of getting in step without any manual adjustment but it may require as many as 10 words (or that number of synchronizing pulses) before it gets into exact phase, on an initial connection or after losing synchronization.

2.1.3 *Transmission Considerations*

Such a signal leads to the following transmission considerations:

1. The signal is rather similar in its spectrum to a telephotograph signal, and as a first order approximation requires about the same type of facility for its transmission.

2. In particular, the signal is expected to require something like the same over-all delay equalization as the telephotograph signal. A brief discussion of this point was given in a paper by one of the authors.⁷ The conclusions there are that the telephotograph equalization limits of ± 0.4 times a signal element duration in envelope delay are generally reasonable, although these are probably overly severe with respect to very fine structure irregularities in the residual envelope delay characteristic. The formal limits which have been set on the envelope delay distortion for the 1,600 bit per second signal are ± 250 microseconds.

3. These limits constitute a rather less severe problem for the bulk of the expected circuits (200–300 miles) than they present for telephotograph circuits which are equalized for 3,000- to 5,000-mile lengths. For the small expected number of very long circuits the problem would of course be the same as for the telephotograph service.

4. The delay equalization problem requires special consideration in view of the nature of practices which have developed in the telephotograph art. The number of circuits which have been equipped for telephotography has in the past been rather small. The adjustment of the delay equalization has involved arithmetical calculation in the process of fitting various manufactured delay equalizer sections to the measured delay distortion. These methods are not generally suitable for large scale operations. It is believed that a process of equalization by prescription will be useable for the 200–300 mile circuits. This is discussed in more detail in Section 3.3.

5. So far, it has not been found expedient to transmit telephotograph signals generally over unmodified N1 carrier or other compandored

circuits, and the same problem is appearing with the proposed data signal. The problem is discussed in more detail in subsequent sections.

6. The specific signal suggested has the disadvantage that multiple levels must be discriminated by the receiver. This increases vulnerability to noise and level changes, as will now be discussed in some more detail.

The specific signal which has been suggested, as noted before, is outlined in Fig. 1. This signal comprises several levels. In the first place it includes a word start indication, on a separate level from the mark and space indications. In the second place the lowest amplitude level (normally a space, but a mark in the "upset" signal, as has been noted) is not made zero, but 0.25 the amplitude of the word start indication or "synchronizing" pulse. The need for this is explained in the discussion on vestigial sideband, below.

Consider for the moment a signal of only two levels; these can be taken as 1 volt and 0 volts respectively. The discrimination between them without error requires that an instantaneous noise pulse at this time be kept to less than $\frac{1}{2}$ volt. In these terms this represents an S/N ratio of 6 db.

The use of 0.25 volt minimum signal means that the amplitude range between maximum and minimum is reduced from 1 volt to $1 - 0.25 = 0.75$ volts. The maximum allowable noise pulse must be $0.75/2 = 0.375$ volts for this signal. This is an 8.8-db S/N ratio, as compared with the previous 6 db. It represents a handicap of approximately 3 db which must be accepted as part of the price of the increased bit transmission rate permitted by the use of vestigial sideband as compared with double sideband transmission.

An additional penalty comes from the use of three as against two signal levels. The spacing signal level is set at 0.625 volt, midway between 1 and 0.25 volt. Discrimination between synchronization and spacing signals can tolerate noise pulses of $(1 - 0.625)/2 = 0.1875$ volt. This amplitude is 15 db below synchronization level. Discrimination between spacing and marking signals tolerates maximum noise pulses of $(0.625 - 0.25)/2 = 0.1875$ volt. This is again 15 db below synchronization level. Thus the signal tolerates a 15-db S/N ratio between synchronization level and the level of maximum noise pulses.

The difference between the approximate 9-db S/N for the two level vestigial sideband signal and the 15 db represents the 6-db handicap caused by the multiple level discrimination in the signal. This is the price paid for a distinctive word start indication.

The price also applies to sudden level changes. In a two level signal

between 1 and 0 volt, a sudden drop of 6 db (without compensating change in the "slicing" or critical level) causes error.

In the two level signal between 1 and 0.25 volt, the permissible drop is reduced to 4 db (ratio of 0.625/1).

In the three level signal, the permissible drop is reduced to 1.8 db (ratio of $(0.625 + 0.1875)/1$).

Automatic gain control and corresponding adjustment of the slicing level ameliorate these conditions to some extent, but the problem is still a serious one; a sudden rise is also serious.

The use of a compandor in the transmission facility exaggerates the situation. Some discussion of the action of a compandor to improve the signal-to-noise ratio for speech is given in Section 3.1. For the moment it can be said that, at the transmitting end, the compandor compresses the range of amplitudes in the speech signal at approximately a syllabic rate. At the receiving end an expansion restores the original amplitude relationships in a complementary manner.

The compression and expansion are matched almost perfectly as far as the human ear is concerned. However data transmission is more vulnerable to short time level irregularities. This allows small imperfections in the amplitude restoration to impose a further penalty in the form of error hazards.

2.1.4 *Other Considerations*

Present ideas call for the acceptance of the signal by the telephone company, and redelivery to the subscriber not in the form indicated by Fig. 1, but in the form of the three separate components listed in 2.1.2. This presents some problems regarding the transmission of such signals over the connecting loops between the subscriber and the telephone office. These are not, however, germane to the general transmission problem and will not be considered further here.

2.2 *Bell Telephone Laboratories and Lincoln Laboratory Vestigial Sideband Systems*

The first of these represents an unpublished exploration, particularly by C. B. H. Feldman and A. C. Norwine, of the possibilities of transmitting moderately high speed data pulses over telephone facilities.

The exploratory system used a carrier at 2,200 cycles; a vestigial band extended up to 2,600 cycles; the nominal effective band extended down to 1,375 cycles; and a roll-off band continued down to 1,100 cycles. In order to reduce the quadrature component resulting from the vestigial sideband the spacing signal was made equal to one-third the marking

signal amplitude. The receiver used AVC both for amplitude control and space-to-mark slicing level adjustment.

The timing of the receiver sampling instants to determine mark and space indications was determined by a flywheel circuit which operated from space-to-mark and mark-to-space transitions in the signal. This is similar in principle to the old Baudot quadruplex arrangements in telegraphy. It was a synchronous system and required a dummy signal for a lining-up period previous to the actual transmission of information. The lining-up was automatic but required some 15 to 50 milliseconds (25 to 80 signal elements).

A number of experiments with the system were made on actual lines. Most of these showed successful transmission, though the error rates were not measured quantitatively. The test showed the signal margin against error on a cathode ray oscilloscope. The wave trace indicated displacements both in the signal amplitude and in the timing of the sampling instant. This margin has been found to correlate reasonably well (in an inverse relationship) with the calculated delay distortions of the circuits used. It also corresponds reasonably well to theoretical expectations.⁷

The system reported on by the Lincoln Laboratory of MIT¹ shows much general similarity to the above. Perhaps the most distinctive feature of difference is in the use of a word start indicating or synchronizing signal, in the form of the high level pulse discussed above and somewhat similar to that used in television for scanning-line synchronization.

2.3 *Double Sideband Systems*

A prototype of these systems has been described by Horton and Vaughan.³ Several models have been derived from the prototype which differ from it, somewhat, in several respects. Large portions of the systems are not germane to the present discussion, and only a brief description will be given of the signals.

An outline of the signal spectrum for the most recent of these derived models is illustrated in Fig. 2. The main signal is handled on a carrier at 1,500 cycles. The bit rate is 750 per second, and the nominal effective band is shown as ± 375 cycles. A schematic roll-off is indicated. The words in this system are of about 100-signal element length, of which 8 are used for synchronization. The synchronization pattern involves a 3-signal element ready pulse on the 600-cycle carrier, simultaneous with the first 3 of the 6-signal element marking pulse on the 1,500-cycle carrier. Following this comes a spacing bit, then another marking bit on the 1,500 cycles only. The next bit is the first information bit.

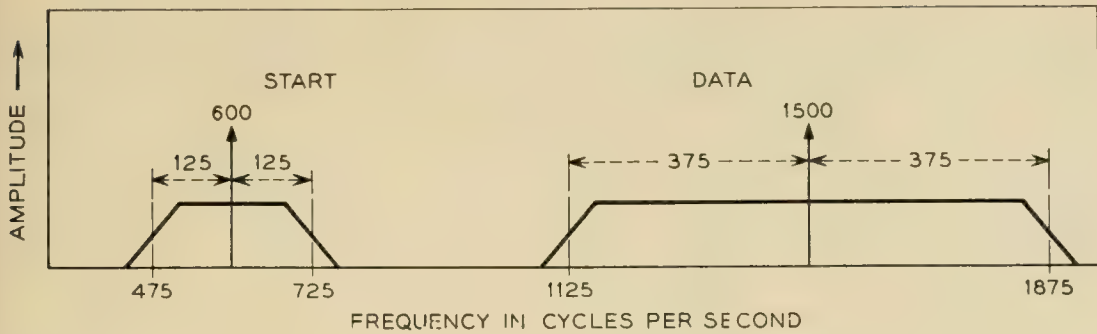


Fig. 2 — Double sideband signal with auxiliary start channel.

Marking consists of maximum carrier amplitude and spacing is zero carrier amplitude. The pulses are shaped both at the transmitting end and at the receiver before sampling.

The earlier prototype system as described by Horton and Vaughan was tested over a variety of telephone facilities (not including N1 carrier) running up to over 12,000 miles in length and found to be quite rugged. For reasons which are to be discussed later, the use of these systems over compandored circuits presents certain extra transmission problems regarding noise and level changes.

2.4 Voice Frequency (VF) Carrier Telegraph

The opposite extreme to the use of the telephone facility as a single band, to carry short duration pulses, is to subdivide the band into a large number of subchannels, each using longer duration pulses. As already noted, this has advantages against impulse noise, and also delay distortion.

There is available for this use the VF telegraph system.⁵ In an AM form (40C1) this subdivides the space into 18 telegraph channels, which can each carry 100 words per minute (or 74 bits per second). A frequency shift form (43A1) is available, to give 17 channels, each to carry 74 bits per second.

The 18 channels use a band of 200 to 3,200 cycles in the telephone facility, and the lowest channel permits only a lower word speed. The 17 channels occupy the band from 350 to 3,200 cycles.

It has been mentioned that the telephone channels which provide the most serious problem for data transmission are those using compandors. The untreated N1 carrier is a principal example of such. Further, the type of circuits which use compandors are apt to be placed in plant which is relatively exposed to impulsive noise.

The principal interest in the telegraph channels, therefore, is to exam-

ine how they fare in application over untreated N1 carrier. There have been some studies of this point. The general conclusion, to be elaborated below, is that although the frequency subdivision (and also the frequency shift) helps against impulsive noise, fewer telegraph channels may be used than over good non-compandored facilities; and at best, the transmission is accompanied by more distortion than expected in a telegraph link of the highest grade. However, a serious possibility for data transmission is indicated.

2.4.1 *Telegraph Tests on N1 Carrier*

Tests on this subject have been carried out by S. I. Cory, J. M. Fraser and others and reported in unpublished memoranda. Before presenting the results of the tests, some background is necessary on the terms in which the results are reported.

The performance is usually evaluated in tests of this kind in terms of a "maximum checkable" telegraph distortion over a short time (about 5 minutes). The "telegraph distortion" represents the displacement from correct timing, of received signal transitions, after the initial mark-to-space transition in the "start" element. These displacements are measured in percentages of the signal element duration. By "maximum checkable" distortion is meant the maximum such displacement that is consistently reproduced in repetitions of the short testing period. This is somewhat larger than the root-mean-square distortion. A larger displacement is, of course, obtainable over a longer testing period. Although this measure of performance has had long use in the telegraph art, other measures are perhaps more readily grasped by and probably of more value to the data transmission engineer. Such a measure, for example, is the error rate.

The error rate may be estimated through the use of the telegraph transmission coefficient.⁸ This is a figure which has been designed by telegraph engineers to indicate the performance of a telegraph circuit, particularly when it is made up of several sections. It is more or less proportional to the square of the distortion, and has the property that, when carefully chosen, it can be added for circuits connected in tandem. A small coefficient thus characterizes good transmission, and correspondingly, a large coefficient characterizes poor transmission.

The correlation between peak distortion over 5-minute intervals and error rate, through the telegraph transmission coefficient, is indicated in Reference 8 and Table II. It is to be understood that the correlation is only a rough one. Particularly at the two extreme ends of the scale, the entries in the table can serve only as a general guide to the performance, and the specific numbers are not to be taken too literally.

TABLE II — CHARACTERIZATION OF TELEGRAPH DISTORTION AND
AND ERROR RATES BY TELEGRAPH TRANSMISSION COEFFICIENT

Distortion		Transmission Coefficient	Errors	
RMS	5 min. peak		1 in n characters	1 in m bits
			(n)	(m)
13.9	30	30	40	2.9×10^2
12.6	27	25	87	6.2×10^2
11.2	24	20	2.5×10^2	1.8×10^3
9.8	21	15	1.5×10^3	1.1×10^4
8.0	17	10	4.4×10^5	3.2×10^5
5.6	12	5	10^9	7.4×10^9
4.3	9	3	10^{12}	7.4×10^{12}
2.5	5	1	—	—

A very brief summary of some of the experiments in the use of VF telegraph over untreated N carrier is given in Table III. This portion of the results covers N1 circuits with compandors which have slightly more noise than the objective which is set for the telephone use of such circuits. The noise was 28 dba at the zero transmission level point, as against an objective of 26 dba at that point. It is noted in Section 3.3.2 that the measurement of noise in these terms is not altogether reliable in the evaluation of its effects on transmission systems that use pulses. Thus, these experiments must be considered as giving only a general indication of the situation.

TABLE III — SUMMARY OF TELEGRAPH PERFORMANCE OVER NOISY
N-I CARRIER LINK

	40Cl (AM)	43Al (FS)	
1. Number of channels.....	6	12	
2. Frequency space used.....	1020	2040 cycles	
3. Words per minute.....	75	75	100
4. Total bits per second.....	342	684	888
Average Channel			
5. Peak distortion.....	16	8	18 per cent
6. Estimated Transmission coeff..	9	2.5	11.5
7. Estimated errors, 1 in.....	10^6	10^{14}	10^5 bits
Worst Channel			
8. Peak distortion.....	25	17	22 per cent
9. Estimated transmission coeff..	21	10	17
10. Estimated errors, 1 in.....	1.5×10^3	4×10^5	5×10^3 bits

The tabulation is first given for an average channel. The performance of the worst channel has, however, also been included to give an indication of its contribution to over-all operation.

Many of the N1 carrier telephone circuits in the plant show lower noise than the objective, and to this extent Table III is somewhat pessimistic. Also some of the tests have shown, particularly for AM, that the performance is somewhat improved by removing the companders. Thus, allowing for both points, better performance can be expected from the average N1 circuit (less than 200 miles) in the plant. The transmission coefficient of 11.5 listed for Item 6 at 100 words per minute might go down to say 9. At 75 words per minute with FS or AM, it might not be over 4.5. It is clear, of course, that with further modifications of the N1 channels, such as to reduce noise exposures, better performance could be obtained.

The broad conclusions that can be derived from these considerations are:

1. The subdivision of the frequency band into telegraph channels, and the use of FS, permit a workable system to be operated over a compandored facility like the N1 carrier without modification. This occurs even when the latter has noise up to and a little over the telephone objective.

2. This workable system under such noisy conditions transmits up to some 350 bits per second with AM, and some 800 bits per second with FS. It is accomplished with an error rate of the order that has been implied for data transmission, even in the worst channel.

3. There is a relatively wide range of performance of the system over different N1 circuits, and the average performance is sensibly better than that under the limiting conditions which have been considered.

2.4.2 *Distribution of Signal in Allocated Bandwidth*

A more extensive discussion of the use of bandwidth is given in Section 3.1, below. However, a few specific points are appropriate here on the band use in telegraph channels.

The spectrum of the original voice frequency telegraph system, based

TABLE IV — USE OF FREQUENCY SPECTRUM IN TELEGRAPH CHANNEL

	AM	FS
1. Channel spacing	170	170 cycles
2. Nominal effective bands	74	74
3. Roll-off band (both sides)	37	26
4. FM swing	0	70
5. Guard band (both sides)	59	0

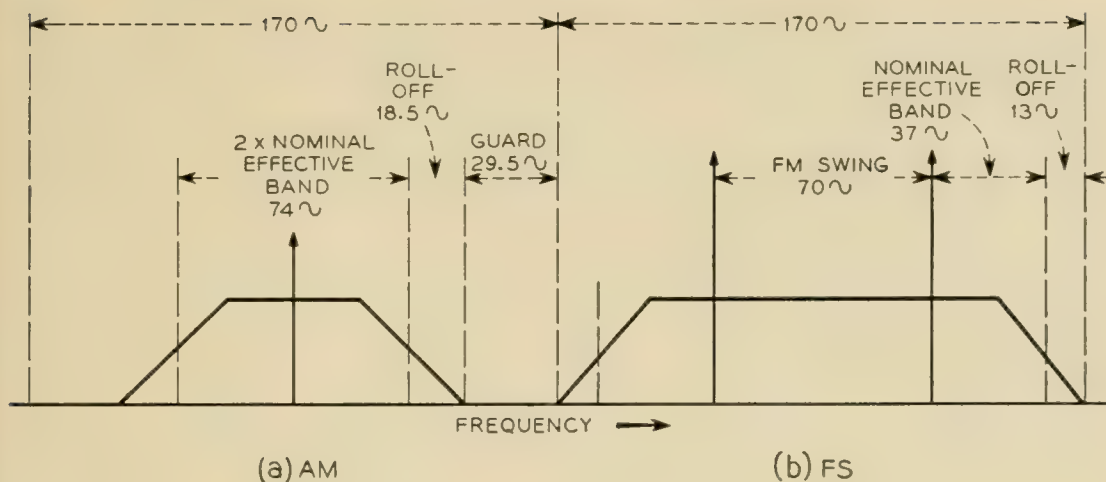


Fig. 3 — Utilization of telegraph channels.

on 170 cycles between carriers, was conservatively developed for the 60 word per minute speed of the time. The 100 word per minute speed has used up some of this conservatism. The use of frequency shift in the same channels has, however, used up the spectrum space even more.

An outline of the band allowances is given in Table IV, and illustrated in Fig. 3. Item 2 of the table is based on the 100 word per minute speed, using double sideband. On this basis, the number is equal to the number of bits per second. This is the minimum double sideband over which that number of bits can be transmitted, according to the Nyquist theory. Each such sideband is sometimes called a "nominal effective band." In practice various allowances are necessary over this minimum.

In the first place a roll-off is necessary because filters are not infinitely sharp, and in addition the nature of the modulation itself forms a roll-off. Roll-off also leads to a signal which is more free of overshoots and generally "cleaner" than when a sharp cutoff is used. Item 3 and Fig. 3(a) and (b) show an allowance for roll-off. For the AM case this amounts to half the nominal effective band. For the FS case there is not quite that much space available.

For the FS signal it is necessary to allow for the frequency swing as Item 4. For the 43A1 system this amounts to 70 cycles. In Fig. 3(b) the spectrum includes the region comprised by the FM swing, and upper and lower sidebands. The upper and lower sidebands as formed by the modulation of a random signal are shown extending respectively above and below the extremities of the swing (instead of only above and below a central carrier, as they would with AM).

A final allowance in Item 5 is a "guard band." This is taken to mean a region in which the signal energy is negligible, but at the same time

a region in which no appreciable interference is tolerated from the adjacent band. The allowance for this is generous in the AM case. In the FS case the roll-off band of Item 3 tends to use up all the space not employed by the nominal effective band and the swing, and nothing is indicated as left for guard band. This is illustrated in Fig. 3(b) by the extremities of the roll-off band reaching the extremities of the 170-cycle spacing. It is to be recognized that the illustrations are diagrammatic. However, the extremities mentioned measure the bands occupied by power from random signal transitions between marking and spacing (as distinguished from mere mark-space reversals), analogous to the bands occupied by power from AM signal transitions. Comparison of Figs. 3(a) and 3(b) suggests how modulated signal components of significant intensity are displaced farther from the edges of the channel with FS than with AM.

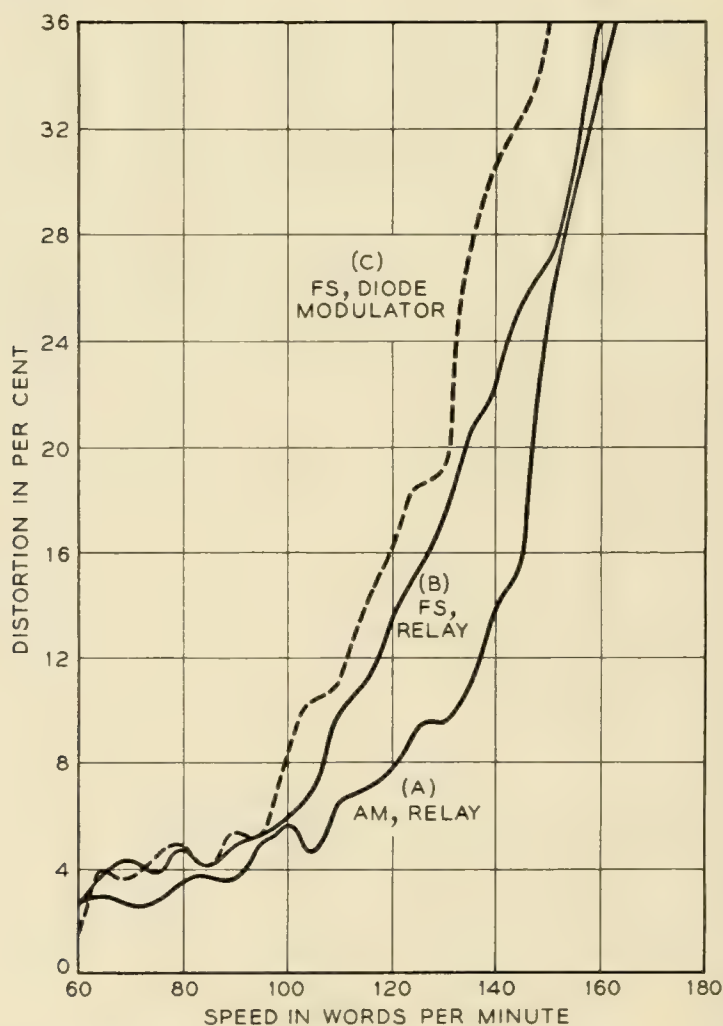


Fig. 4 — Experimental relation between telegraph distortion and speed, with fixed channel filters (after Jones and Pfeiffer⁹).

The conclusion is reached that there is hardly any excess conservatism in the 43A1 system. Some confirmation of this is indicated in a paper by Jones and Pfeleger.⁹ Fig. 4 reproduces some curves presented in that paper. The curve at (B) (from Fig. 5 of that paper) shows that the FS telegraph rises rapidly in distortion above the 100 word/min speed. The curve at (A) (taken from the same Fig. 5) for level-compensated AM shows a substantially broader and somewhat lower curve in this region. From curve (C) (taken from Fig. 3 of the paper) for diode modulated signals, it is seen that the sharp rise in distortion for FS is even more accentuated than for the relay modulator. This indicates how much more characteristic distortion FS exhibits than AM because of the sharper roll-off of its signal bands within the confines of the same 170-cycle channel spacings. Of course, as the word speed is raised beyond the practically usable values, either type of modulation leads to so much power in the filter cutoff regions that the characteristic distortions become more or less indistinguishable.

2.5 "Polytonic" System

This is a frequency discrimination system experimentally proposed for toll and local signaling.⁴ It works on 5 channels at speeds of 100 and 300 decimal digits per second (or the equivalent of some 330 to slightly under 1,000 bits per second). It has some similarities to an earlier multi-frequency system,¹⁰ but is faster.

The distinguishing characteristic of the polytonic system lies in the mathematical theory which has been followed to reduce interchannel interference. This analysis makes use of the theory of orthogonal functions, and is similar to that used in the computation of Fourier components. The mathematical analysis leads to an ideal receiver design for minimum interchannel interference. The ideal detector in this receiver closely resembles the conventional homodyne detector. The detector actually used, however, represents a practical simplification of the latter. A complete description cannot be given here, but it may be noted that the theory leads to a need for synchronization of the signal elements in the five channels, and to the setting of an exact timing instant for the sampling of the received wave to obtain minimum interchannel interference. The indication for this instant is obtained from the use of a sharp wave-front pulse in the marking channels.

Tests with the 100 decimal digit per second system (100 decimal digits normally correspond to 332 binary digits) indicated that it gave

good operation⁴ over toll circuits of comparatively limited length (350 miles for four-wire voice frequency, 1,900 miles for K carrier). The higher speed system was designed only for local plant.

It is clear that this system has too low a bit rate for the application contemplated, even in the faster form. It is also not generally adaptable to the variety of circuit lengths which are expected to be encountered.

III. UTILIZATION OF THE TELEPHONE CHANNEL

The discussion herewith covers a broad examination of some major characteristics of telephone communications facilities, to evaluate their bearing on the choice of a system for the data transmission service outlined before. It is of course clear that different conclusions might be reached for other types of service.

The first item is an outline review of the different types of message telephone facilities in the plant. This is followed by an analysis of the different possibilities in the use of the frequency spectrum, of noise, and of delay distortion, in the application of the data signals.

3.1 *Telephone Facilities*

There is a rather wide variety of facilities to be found in the telephone plant, to be examined with respect to the factors that are germane to the present question.

The first of these factors is the frequency bandwidth capability of the facility. For message-type voice circuits, this is generally characterized as being three kilocycles (with the exception of "emergency banks," which are substantially narrower, and are not to be considered as useable for data transmission).¹⁷ However, some of the telephone circuits in the plant, aside from the emergency banks, are also somewhat narrower than 3-kc, and in any case, not all the band is effectively usable for data transmission. As will be seen, the net available band is, in practice, about half of the 3kc.

Part of the reason that not all of the frequency band is effectively usable is that the circuit shows delay distortion. This tends to become large at both the lower and the upper edges of the band. Some details of the delay correction are discussed further below.

Another impairing factor in telephone facilities is the nonlinear distortion encountered. In voice frequency facilities this comes from amplifiers and loading coils, and increases progressively with circuit length. In carrier facilities the nonlinear distortion arises almost exclusively at

carrier terminals, in the part of the circuit where the signal is at voice frequency. In such a case, the distortion increases with the number of times in the telephone facility that the signal is modulated down to voice frequency. Second order nonlinear distortion tends to develop modulation products in the lower portion of the transmission band which are a source of potential interference with the signal.

Another impairment encountered in carrier telephone channels is a slight frequency shift; that is, a 1,000 cycle input may appear at the output, say at 998 cycles. This occurs because modulator and demodulator frequencies are not identical. With independent oscillators on recent systems this shift may amount to some two cycles. With older systems it can run from 5 to 10 times as much. The effects of this shift are discussed below. The frequency shift may be avoided by working double or vestigial sideband and using an envelope detector, or in some carrier systems by locking the oscillators in a constant frequency network. This locking may or may not result in close phase synchronization of the carriers, depending on the method used to lock and the particular carrier system involved.

Still another factor is the use, or not, of "compandors." A compandor compresses the range of speech volume in the impressed line signal and correspondingly expands this range at the receiver. This raises the line signal level during periods of low speech power, and lowers it during periods of high speech power without, in principle, affecting the final received level. The effect is to reduce the final noise in periods of low speech power, and increase it during periods of high speech power. A listener is less perceptive to the noise during high speech power levels than low. By this means, it has been found that the telephone circuit can be engineered to some 23 db more noise (and also crosstalk and similar forms of interference) than it can without the use of a compandor.

In the case of data signals, however, the influence of noise in causing error is not very much different whether the signal is marking or spacing. Thus, there is no "compandor advantage"* (indeed there is a certain disadvantage as pointed out earlier), and facilities that have been engineered to be entirely satisfactory for voice transmission are effectively some 23 db more noisy for data transmission. As a practical matter it appears desirable to remove compandors from circuits used exclusively for data.

A short listing is presented here of the various types of message facili-

* Perhaps a simpler way to think of it is that all possible "compandor advantage" has already been obtained in a data system by using the best combination of amplitudes for mark, space, start, etc.

ties most frequently found in the telephone plant, and some comments are made on each.

3.1.1 *Voice Frequency Circuits*

There is a variety of open-wire facilities of this type. They are mostly short, and two-wire. Thus, repeaters can to advantage be turned one-way for data service. Delay correction is discussed later.

Voice frequency cable facilities over more than a very short distance are loaded. This gives appreciable delay distortion. The loading used is indicated by a letter denoting the spacing, followed by a number denoting the loading coil inductance. Thus "H-44" means 6,000-foot spacing, of 44-millihenry coils, and "B-88", 3,000-foot spacing of 88-millihenry coils. Conductor capacities range from 0.62 microfarads per mile for toll circuits, to 0.82 microfarads per mile or sometimes even higher for local circuits. This affects the delay distortion.

3.1.2 *Type-C Carrier Circuits*¹⁶

This is an open-wire three-channel system operating at different frequencies in opposite directions, over the same pair. Historically there has been a variety of C systems developed, but only the C-5 system exists in any extensive quantity. The upper frequency cutoff in the voice channel is well under 3 kc. The delay distortion varies widely with the specific channel and direction of transmission. There is a variety of channel frequency allocations, and the distortion varies with this also. The delay distortion over some channels increases rapidly above 2,400 cycles. The frequency shift discussed before may be as much as ± 20 cycles.

3.1.3 *Type-N Carrier Circuits*⁶

This is a short-haul twelve-channel system for use over cables. Because of its economy it has been extensively introduced. Its principal characteristic, in the application of data circuits, is that it uses companders. It therefore presents a noise problem. The delay distortion, introduced almost exclusively by the terminals is not excessive, and depends very little upon circuit length between the terminals. The N system uses double sideband transmission, and therefore exhibits no frequency shift between input and output signals.

3.1.4 *Broadband Carrier Systems Using A Channel Banks*¹⁵

There is a variety of carrier systems designed for paired cable, coaxial cable, open wire, and radio, that use a standard grouping of twelve channels with associated filters, known as an "A channel bank." The delay distortion in these associated filters constitutes nearly all of the distortion measurable over the complete system. These are single sideband systems, and unless the local modulator and demodulator carrier

supplies are locked in a constant frequency network, frequency shifts of some ± 2 cycles may be expected between input and output.

For paired cables, these are known as K1 and K2 systems. For coaxial, they are L1 and L3, for open-wire, J, and for microwave radio, TD-2.

3.1.5 *Other Broadband Carrier Systems*

An O carrier system has been developed for open wire, and combinations of it are used with N for open-wire and carrier. These are compandored systems.

3.2 *Use of Bandwidth*

This section examines the more important factors which affect the choice of how the available bandwidth of a facility is to be used, either in one band or a subdivided band.

3.2.1 *Baseband Transmission*

This is the simplest type of transmission. It is used in telegraph loops and other short distance telegraph transmission. A mark is indicated by placing marking voltage across the wire line, and a space by placing spacing voltage. In the simplest systems the latter is zero. In "polar" systems it is the negative of marking voltage.

The frequency spectrum of the signal runs down to and includes dc, as illustrated by the solid lines in (a) of Fig. 5.

With many transmission facilities it is difficult or impossible to transmit the dc; i.e., the circuit cuts off as is illustrated by the dotted lines. In such cases it is impossible to distinguish between a permanent mark and a permanent space.

Extra pulses can, however, be added to the signal to insure that marks or spaces are not permanent, but are relieved by the opposite signal in some maximum interval of time. In such cases the received signals can be clamped on mark or space signals and the opposite condition can be readily distinguished. This is sometimes called "dc restoration," and strictly speaking the system ceases to use baseband transmission. It may be designated as "modified baseband transmission." Methods other than clamping have been suggested for dc restoration.

Reverse pulses can be systematically inserted after each mark or space pulse, according to various patterns.¹¹ Two suggested are "dipulse" and "dicode" pulses. Such signals approach carrier signals, which are discussed below.

The principal weakness of baseband transmission appears when it is sent over C carrier or other single sideband telephone facilities, where the recovered signal may vary in frequency from that sent. This causes a distortion of the received pulse which makes it difficult to recognize.

An analysis of this point is given in Appendix I. It is concluded that while there may be long range possibilities in baseband transmission, it requires more study, and it will not be considered further in this paper.

3.2.2 AM Carrier Transmission

The simplest of this type is double sideband transmission, as illustrated by the full line of Fig. 5(b).

A comparison of the susceptibility to noise of this arrangement, with that of baseband transmission, is considered in the next section.

A further consideration required is susceptibility to nonlinearity in the facility. Second order modulation leads among other things to a rectification of the signal back to baseband. This is indicated by the dotted lines in Fig. 5(b). After such rectification of the signal by the facility, it is impossible to separate any overlapping portions of the signal between the baseband and lower sideband. Some overlap is shown. This interference was first considered in telephotography¹² and is known as "Kendall effect."

The possibility of Kendall effect may be eliminated insofar as second order modulation is concerned by moving the carrier frequency high

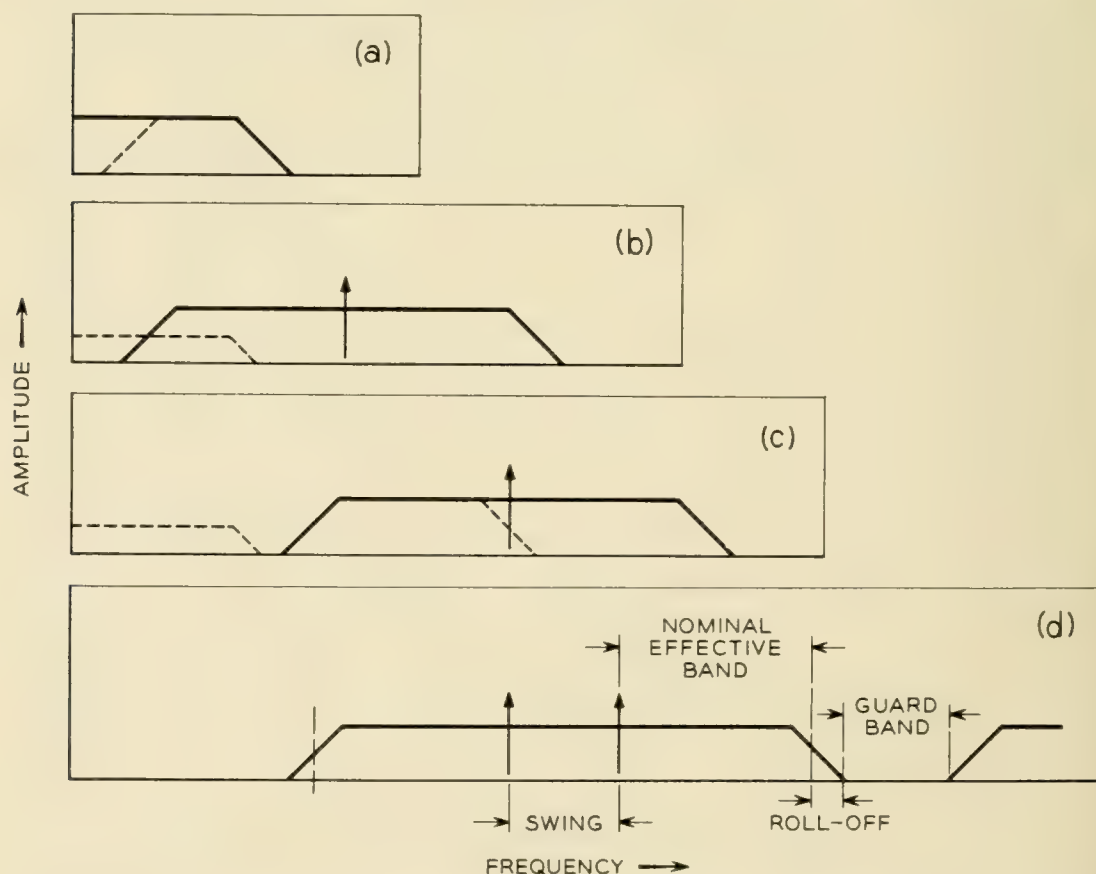


Fig. 5 — Spectra of signals with various modulations.

enough to prevent such an overlap. This is indicated in Fig. 5(c). It does not prevent third order modulation effects.

It has been found necessary to allocate frequency bands thus to avoid the overlap discussed in the transmission of high grade telephotography over telephone type facilities. It has also been noted that allocation with such overlap was undesirable in some data transmission experience. However, the question has not been resolved in complete detail. For the data service under consideration, which is expected to show a very low error rate, it is deemed conservative to allocate bands without the overlap.

This conclusion then leads to wasting a certain part of the lower frequency range. It is still possible, however, to use this range for an auxiliary signal, as in the system illustrated in Fig. 2, if the auxiliary signal occurs only during word starts.

The double sideband frequency range, as indicated in Fig. 5(b) or 5(c), is about twice the baseband range of Fig. 5(a). It is possible to reduce this extension by cutting down one of the sidebands to a "vestige" of itself and sending carrier at a reduced level, as indicated by the diagonal dotted line about the carrier in Fig. 5(c). This was proposed by Nyquist.¹³ It is done at the expense of an increased vulnerability to noise, which in total amounts to some 5 or 6 db in certain typical cases.^{7, 11} In Section 2.1.3 discussion was given to account for 3 db of this. In the references cited herewith it is noted that vestigial sideband transmission is accompanied by an interfering component (called a "quadrature component") which accounts for the other 2 or 3 db.

3.2.3 *FM Carrier Transmission*

Certain additional immunity to noise is gained by the use of frequency modulation (or "frequency shift") of the carrier. The immunity which can be obtained against impulse noise can be even greater than that against random noise, provided that the receiver is precisely tuned. This was noted in Section 2.4 in connection with voice frequency telegraph.

The noise immunity obtained from FS is in part due to the use of a higher average power and is at the expense of a wider frequency band as illustrated in Fig. 5(d). In addition to all the band that is used for a double sideband system, a space must be allowed for the swing. FS is also much less vulnerable to sudden level changes than DSB and thus may well be preferable to DSB for medium bit-rate service. As shown in Table I, these advantages can be obtained with only a small sacrifice of bit rate compared to DSB for equal bandwidths.

So far the single channel broadband FS system has not been generally

used over wire circuits. Hence its exact performance, particularly with impulse noise, is only estimated.

On occasion not all of the band illustrated in Fig. 5(d) is allowed for. This leads to an increase in distortion of the signal, which has some similarity to a very close-in echo. It uses up some of the additional noise immunity provided by the FM, as an engineering compromise.

Another direction along which the FM system may be practically extended is to use four instead of two (marking and spacing) frequencies. This would double the bit capacity at the expense of only a moderate widening of the frequency band and somewhat tighter requirements on noise and delay distortion (but not of level regulation, which would be required for a similar extension of the AM signal).

3.2.4 *Multichannel Systems*

It is possible to divide up the entire frequency band available into a number of separate channels and use any one of the various carrier systems which have been described, in each individual channel. This may be done because the nature of the information transmitted may be better adapted to the narrow channel, as in conventional telegraph. It permits certain elements of flexibility in layout, and offers certain noise advantages (and also disadvantages) as discussed in the next section.

In an idealized way one can proportion the various allowances for nominal effective band, roll-off, guard band, and swing (FS) in the same proportion in which they would occur in a single broad channel over the whole facility. Thus no frequency space would be lost by the subdivision. In practice, however, subdivision usually does lead to some actual loss in the frequency space.

A significant limitation to frequency subdivision lies in nonlinearity of the facility. This leads to modulation products between the various channels, which interfere with other channels.

In the case of voice frequency telegraph and other multiple channel systems the modulation effects are mitigated by allocating the carriers at odd multiples of a basic frequency. That is, any given carrier f is set at $f = nk$, where n is odd and k is a basic figure. Then the three second order modulation products are

$$2f = 2nk,$$

$$f_1 + f_2 = (n_1 + n_2) k,$$

$$f_1 - f_2 = (n_1 - n_2) k.$$

TABLE V — USE OF TELEPHONE BAND BY VARIOUS DATA SYSTEMS

System	Modulation	Max. No. of Channels	Bits/Sec. per Channel	Total Bits/Sec.
1. Proposed	VSB	1	1600	1600
2. Exploratory	VSB	1	1650	1650
3. Exploratory	DSB	1	750	750
4. Polytonic {	toll	5	100	500*
	local	5	300	1500*
5. 40C1 Teg.	DSB	18†	74††	1332††
6. 43A1 Teg.	FS	17†	74††	1258††

* Not realizable with 2 out of 5 codes used.
† Not realizable over some facilities.
†† Based on 100 word per minute channel capability.

All these products are necessarily even multiples of k and therefore always fall half way between carriers. This allocation, however does not permit mitigation of third order modulation effects.

3.2.5 Experience

Table V reviews systems which have been discussed earlier in this paper, to indicate the extent to which they use a general telephone channel facility, in terms of the bit rate output. It is clear, of course, that the various systems are not engineered to the same conservatism. These differences have already been commented upon.

The general conclusion which one can reach from the table is that the use of the medium in a single channel gives possibilities of a higher bit rate than subdivision. However, it is to be kept in mind that the telegraph facilities are conservatively engineered. Further as noted, they can be used to the full extent indicated only over the broader band telephone channels. For example, the full 18 telegraph channels can not be used over a C carrier telephone channel.

3.3 Noise

A general theory regarding the influence of noise on digital systems was presented in 1948 by Oliver, Pierce and Shannon.¹⁴ This was discussed further at a symposium.⁷ Several additional points are considered here, relating to the effect of noise on a data transmission service of the type contemplated.

3.3.1 Effect of Channel Subdivision on Vulnerability to Noise

This has been suggested earlier at several points. A more detailed discussion is given of the effects in Appendix II.

The conclusions reached there are, broadly:

- 1. Channel subdivision has comparatively small effect on vulnera-

bility to random noise. The small effect which does occur is a disadvantage which can run up to some 3 or 4 db for ten or twelve channels, for the multichannel as compared to the single channel system.

2. Channel subdivision is advantageous over single channel use, with regard to vulnerability to impulse noise.

3. Channel subdivision is disadvantageous compared to single channel use, with regard to vulnerability to single frequency noise.

3.3.2 *Noise Effects in Vestigial Sideband System*

Some brief discussion of the noise problems in the vestigial sideband system under consideration has already been presented in Sections 2.1.3 and 3.1. That such problems may be important in the use of actual telephone facilities has generally been checked by some unpublished tests carried out by J. Mallett, of Bell Telephone Laboratories.

The noise problem is most serious in the application of data transmission over N1 carrier. It is particularly significant for the N1 installations previous to the most recent. The most recent installations are engineered with distinctly more conservatism in regard to noise performance. As is noted in Section 3.1, the principal characteristic of the N1 system that affects this application is that it uses a compandor and that its design for telephone use assumes a reduction of the noise by this compandor. The reduction then is not realized for data signals. A second characteristic is that N1 channels are exposed to impulse noise. The channels are of course designed to limit such noise to the extent that it will not sensibly impair telephone speech. But short data pulses are more vulnerable to the impulse noise than speech. The 2B noise meter, which is normally used for telephone noise measurements, does not read sharp noise impulses according to their effect on data signals, and other methods of measurement have been explored.

A summary of some of Mallett's results is plotted in Fig. 6. Here the reading on the 2B meter is compared with that on a level distribution recorder which records peaks of 1 millisecond or longer. The "one per cent" point is noted, which means that one per cent of the one second intervals in the period of measurement contained one or more peaks (of 1 millisecond or longer duration) of the amplitude indicated. This is found convenient as a measure of the error frequency tolerated in the system considered (one in 10^5 bits).

The heavy dots indicate the correlation between the two sets of noise readings on idle tested channels of an operating N1 system (other channels in the system being busy due to normal use). All channels had compandors in. Dots that mark the boundaries of a group of tests (usually a single system in the group) are connected by straight lines. The open

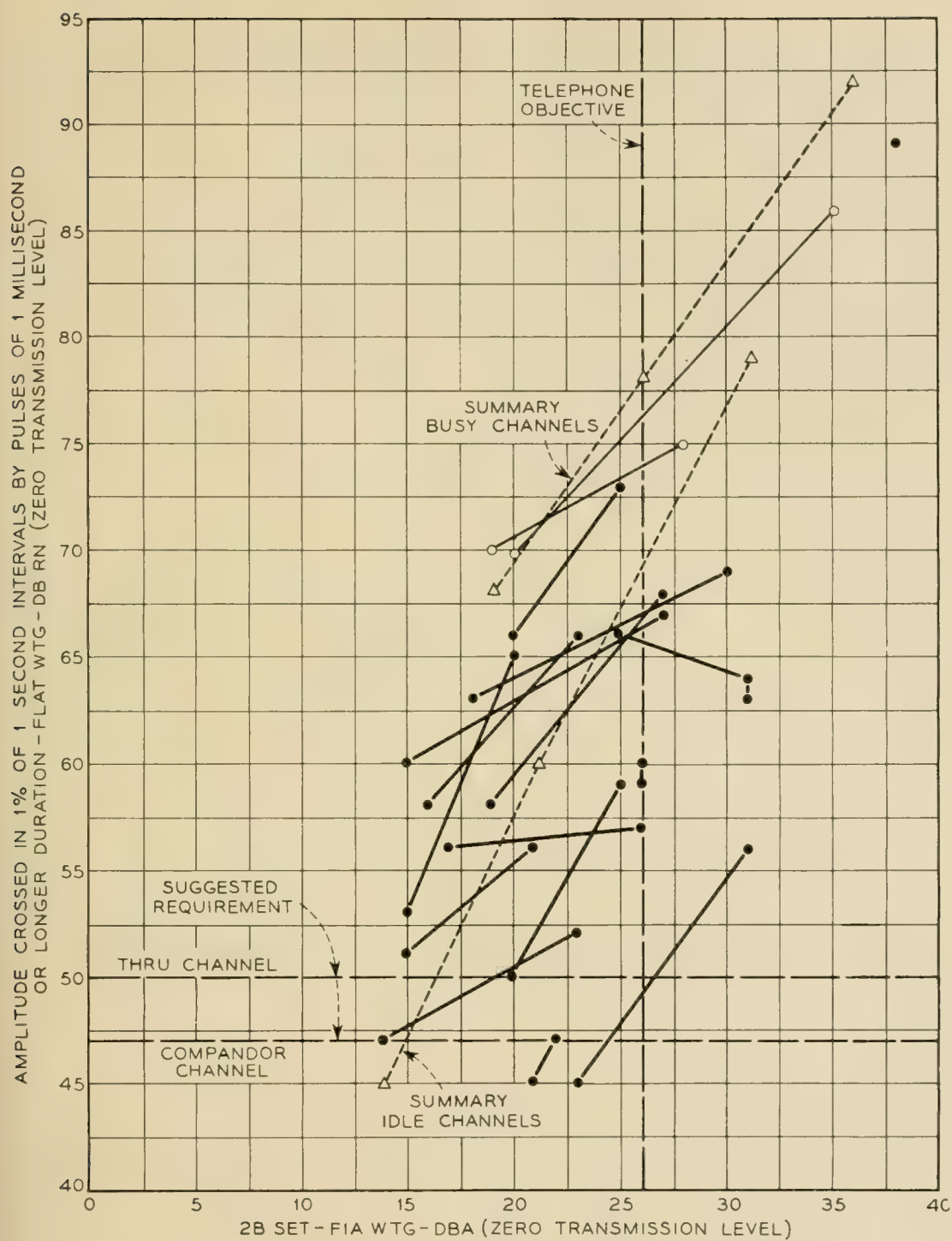


Fig. 6 — Noise measurements in N1 carrier facilities.

dots refer to an estimate (computed from the known properties of the compandor) of what the noise peak levels would have been, had the channel tested been busy with an operating data transmission signal (of the general type illustrated in Fig. 1). Under this condition the compandor setting would of course be different from that of an idle channel.

In such cases the 2B meter reading is also adjusted to the effective message circuit noise which would have existed in the presence of speech on the tested channel.

Open triangles connected by fine dotted lines indicate summary plots for the idle channels, and for the busy channels.

Examination of the plot shows that there is only a general correlation between the 2B set and level distribution recorder readings, as summarized by the dotted lines. Thus the 2B meter is not a reliable instrument to denote how noisy a circuit is for data transmission of the speed used in the proposed project.

The telephone objective for N1 circuits, as read on the 2B meter, is indicated by the vertical dotted line. This shows that most of the circuits (actually 55 out of 62) met the objective.

The suggested requirements for data are shown by the horizontal dotted lines. The "through" or noncompandored channel has a 3 db more lenient requirement (as indicated during some of the tests) than the compandored one. This results from the penalty mentioned earlier which compandors impose on in a data circuit, as compared with the 23 db or so advantage that it introduces for voice transmission. Only a few of the circuits measured (actually 10 out of 62) met the suggested requirements.

The principal conclusion reached from these measurements is that where used for a data service of the type considered, with vestigial sideband transmission, most of the hitherto installed N1 circuits (and probably other compandored circuits) will require modification to reduce noise exposure. Also noise measurement will be more complicated for such a service than for telephony.

3.4 *Envelope Delay Distortion*

Some simple theoretical considerations⁷ have shown that the envelope delay distortion limits for a telephotograph circuit generally also hold for a data transmission circuit of the same speed. The principal difference is that less emphasis need be placed for data circuits upon fine structure deviations of the envelope delay as plotted on a frequency scale. In general, distortion of ± 0.4 signal element in the important part of the band has been found to give a signal-to-noise impairment of some 3 db in signal reception. This has been assumed here as a tentative engineering objective.

In accordance with this, the envelope delay requirements for service with the vestigial sideband signal consideration have been set at not to exceed 500 microseconds (± 250 microseconds) between 1,000 and 2,500

cycles. This contains an element of conservatism inasmuch as the strict requirement is really fully implied only on the nominal effective band (1,200–2,000 cycles). The signal power is reduced in the roll-off and vestigial bands, respectively 1,000–1,200 and 2,000–2,500 cycles, and some corresponding liberality may be expected there.

The delay distortion constitutes a more serious problem with a faster system as compared with a slower one, in part because of the wider frequency band occupied by it, and in part because 0.4 signal element represents a more severe tolerance in microseconds for a shorter element than for a longer one. Consequently, the limits given represent about as severe tolerances as may be expected to be needed with the use of a telephone channel.

The distortions of various circuits have been considered to estimate the order of the problem involved in meeting the proposed requirements over links of 100 to 500 miles.

The following conclusions are reached first for the vestigial sideband signal, and after this for the slower systems.

3.4.1 *Facilities Requiring No Treatment*

As already noted, K2, L1, and L3 carrier, and TD-2 microwave, use "A" channel banks to separate the individual channels, and these give the dominant delay distortion. This amounts to a maximum of about 200, and a minimum of 150 microseconds, according to the exact combination of filters used. This figure is for one link of transmitter and receiver. A single section delay equalizer can cut the maximum residual to about 80 microseconds. It is concluded that these facilities present no important delay distortion problems. An N1 carrier link gives a maximum delay distortion of 220 microseconds, which can be reduced to 50 microseconds by one section of equalizer. This, then, also presents no serious problem.

3.4.2 *Facilities Treated by Simple Prescription*

The delay distortion of H-44 voice frequency cable in the 1,000- to 2,400-cycle range runs to slightly under 900 microseconds for 300 miles, if the cable is of standard toll capacitance (0.062 mf per mile), and to slightly under 2,000 microseconds if of higher local plant capacitance (0.084 mf per mile). The use of about one section of equalization per 100 miles reduces the residual to less than 330 microseconds for the low capacitance cable. For the higher capacitance cable, about three sections are needed per 110 miles. The J-2 carrier uses A channel banks, but has, in addition, directional separation filters at each repeater. This gives maximum and minimum distortions, respectively, for 100 miles, of slightly under 300 and slightly under 160 microseconds. The precise

distortion in any given channel depends upon its proximity to the cut-off of the directional filter. For 500 miles the figures are slightly over 500 and slightly under 50 microseconds. With the same single section of equalization the maximum figures are reduced to about 100 microseconds for 100 miles, and about 300 microseconds for 500 miles. To carry out this equalization requires only rudimentary information on the general nature and correction of delay distortion. If moderate care is used in the prescription of equalization on a packaged basis no delay measurement of the circuit would in general be needed, though it is recognized that some difficult cases may arise.

3.4.3 *Facilities Requiring More Involved Prescription*

The delay distortion in C-5 carrier¹⁶ is influenced to a dominating extent both by channel and directional separation filters. It varies in a complex fashion from channel to channel, and according to the direction of transmission. Its correction thus requires more involved prescription than is required for the other types of circuit. In some few cases measurement may be necessary. The distortion of H-88 voice-frequency cable runs from some 1,400 to over 3,000 microseconds per 100 miles according to capacitance. For 20 miles of H-174 toll cable the distortion is slightly under 1,400 microseconds, and its use is not contemplated.

3.4.4 *Data Systems Requiring No Corrections*

The delay distortion problem is practically non-existent for the slower systems. For the double sideband systems some delay correction may be needed if long heavily loaded circuits are used or perhaps for some other rare unfavorable situations, but otherwise no correction is necessary. No correction is needed for the telegraph systems.

APPENDIX I — BASEBAND SIGNAL DISTORTION CAUSED BY CARRIER FREQUENCY SHIFT

A simple analysis of the phenomenon may be considered. Let the voltage input, as in Fig. 7(a), be a raised cosine pulse between the angular arguments of $-\pi$ and $+\pi$. That is

$$V_i = 1 + \cos \Omega t, \quad (1)$$

where Ω/π is the envelope frequency. When this is transmitted on the carrier, $\cos \omega t$, the carrier signal voltage is

$$V_c = (1 + \cos \Omega t) \cos \omega t, \quad (2)$$

$$= \cos \omega t + \frac{1}{2} \cos (\omega - \Omega)t + \frac{1}{2} \cos (\omega + \Omega)t. \quad (3)$$

When the carrier and one sideband are removed (say the lower side-

band), V_c becomes (neglecting the factor $\frac{1}{2}$)

$$V_c = \cos (\omega + \Omega)t. \quad (4)$$

At the receiving end, V_c is modulated with a carrier which may momentarily differ in phase from the signal carrier by angle φ , giving

$$V_o = \cos (\omega + \Omega)t \cos (\omega t + \varphi), \quad (5)$$

$$= \frac{1}{2} \cos [(\omega + \Omega)t - (\omega t + \varphi)] + \frac{1}{2} \cos [(\omega + \Omega)t + (\omega t + \varphi)]. \quad (6)$$

The lower frequency part of V_o constitutes the recovered signal, and it is extracted by a filter that attenuates the higher frequency part. Thus, again neglecting the factor of $\frac{1}{2}$,

$$\begin{aligned} V_r &= \cos (\omega t + \Omega t - \omega t - \varphi), \\ V_r &= \cos (\Omega t - \varphi). \end{aligned} \quad (7)$$

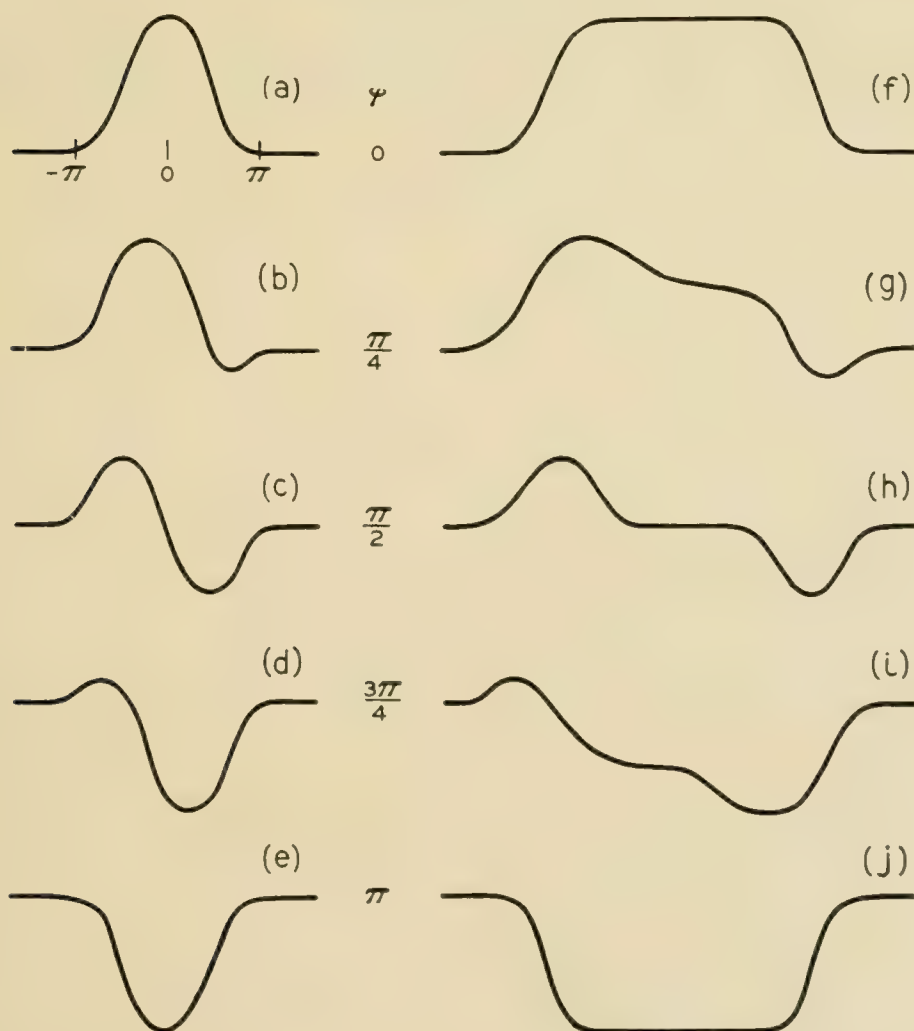


Fig. 7 — Distortion of baseband pulse signal in single sideband carrier facility.

The angle φ progresses through 0 to 2π per cycle of difference between the original and recovered frequencies. That is, if this difference is 2 cycles, then φ progresses from 0 to 2π twice each second.

The effect of the progression from 0 to π is illustrated in Figs. 7(a) to 7(e). When the phase is equal to π , as at 7(e), the signal is "upset"; i.e., marks are changed to spaces, and vice versa. When the phase is equal to $\pi/2$, the signal is effectively differentiated, as at 7(c).

In general, the signal as specified in (1) includes harmonics of Ω . An illustration of such a signal several dots long is shown in Fig. 7(f). As before, when $\varphi = \pi$, the complete signal is upset, as illustrated in Fig. 7(j). When $\varphi = \pi/2$, as at 7(h), the signal is more or less differentiated. The differentiation is not exact because the successive harmonics are not weighted according to order, as in true differentiation.

The distortions caused by the progression of φ are what make it difficult to recognize the signal at the receiving point. Some suggestions have been made for correcting the indication when the signal is upset, as in Figs. 7(e) and 7(j). It is more difficult, however, to take care of the intermediate cases, particularly 7(c) and 7(h).

APPENDIX II — EFFECT OF CHANNEL SUBDIVISION ON VULNERABILITY TO NOISE

There are three broad categories of noise to which the system may be exposed. These are:

1. Random noise which is not localized in time nor frequency.
2. Impulse noise which is highly localized in time but covers a broad frequency spectrum.
3. Single-frequency noise, which is highly localized in frequency but which lasts a significant time (or substantial number of bits).

We can assume two systems, A used as a single-frequency band, and B divided into ten channels. Correspondingly, therefore, signal pulses of unit duration over system A, are of 10 units duration in each channel of system B.

Random noise having uniform spectral distribution, and power W in each channel of system B, cumulates to power $10W$ in system A. Signal power P in each channel of system B cumulates to $10P$ for the total system. If signal power $10P$ is used in system A, the two systems are at a stand-off in signal-to-noise ratio for this type of noise.

In practice, the power capacity to handle the signals for system B must be made a few db higher than indicated by $10P$ to allow for occasional peaking caused during instants of unfavorable phasing among the vari-

ous channels. Thus, the multichannel system B is really worse off, by that small amount, than system A. This may amount to some 3 or 4 db in ten or twelve channels.

There are occasions where crosstalk into other facilities sets the level permitted for the signal. It may be that concentration of the signal onto a single carrier aggravates this interference, as compared with that from a multi-carrier signal. In such cases system A may be penalized by a few db, as compared with system B.

Impulse noise shows correlation among the phases of its spectral components. Thus a noise pulse of voltage amplitude N in each channel of system B, cumulates to voltage amplitude $10N$ in system A or 20 db greater. On the other hand, a signal of amplitude S in each channel of system B, cumulates to an RMS voltage amplitude $\sqrt{10} S$, or 10 db greater, over the total system (plus a peaking correction which may be positive or negative as just mentioned). Thus, the single channel A system is at a disadvantage of 10 db, less the peaking adjustment, with respect to the B system.

A further adjustment may be needed, because a single noise peak that affects all 10 channels of B, or 10 bits, may affect only one bit of A. This adjustment depends upon the word grouping of bits which is used. It may be neglected if all the 10 bits of B are in the same word, and if, in one word, an error of 10 bits is effectively no worse than an error of 1 bit.

Single-frequency noise lies at the opposite extreme of the gamut from impulse noise. The vulnerability of a signal pulse to single-frequency noise varies according to the relationship between the frequencies of the noise and of the signal carrier in the utilized signal band.

The pattern of sensitivity to noise over an individual channel can be expected to be about the same for a narrow band as for a wide-band channel. Thus the pattern of sensitivities in the single channel of system A is repeated in each channel of system B on a 10 times finer frequency scale. The required S/N ratio in any one channel of B remains the same as that for A.

In system B each of the ten channels must put out only one tenth of the power of the single channel of system A (less the correction for peaking which as before may be positive or negative). Thus any one of the ten channels of B is 10 db (plus a peaking correction) more vulnerable to single frequency noise than the single channel of A.

It must be noted that there are occasional special circumstances where the single frequency noise may be persistent and steady. The multichannel system B may in such cases have an advantage in permitting

the one channel affected to be dropped, and the others to be worked entirely free from this interference. This of course reduces the total bit rate.

To summarize the discussion in a general philosophical way, it can be said that there is advantage in multiplexing the signal in the manner that makes it as different as possible from the type of noise to which it is expected to be the most exposed. If the predominant noise is in short duration pulses, the most advantageous signal is in long duration pulses with frequency discrimination multiplex. If the noise is in longer duration single frequencies, the most advantageous signal is in very short pulses with time discrimination multiplex.

VI. REFERENCES

1. J. V. Harrington, P. Rosen, D. A. Spaeth, Some Results on the Transmission of Pulses Over Lines, *Symposium on Information Networks*, **III**, pp. 115-130, Polytechnic Institute of Brooklyn, April, 1954.
2. F. W. Reynolds, A New Telephotograph System, *B.S.T.J.*, **15**, pp. 549-475, October, 1936.
3. A. W. Horton and H. E. Vaughan, Transmission of Digital Information Over Telephone Circuits, *B.S.T.J.*, **34**, pp. 511-528, May, 1955.
4. C. A. Lovell, J. H. McGuigan and O. J. Murphy, An Experimental Polytonic Signaling System, *B.S.T.J.*, **34**, pp. 783-806, July, 1955.
5. A. L. Matte, Advances in Carrier Telegraph Transmission, *B.S.T.J.*, **19**, pp. 161-208, April, 1940; J. R. Davey and A. L. Matte, Frequency Shift Telegraphy — Radio and Wire Applications, *B.S.T.J.*, **27**, pp. 265-304, April 1948.
6. R. S. Caruthers, The Type N1 Carrier Telephone System; Objectives and Transmission Features, *B.S.T.J.*, **30**, pp. 1-32, January, 1951.
7. P. Mertz, Transmission Line Characteristics and Effects on Pulse Transmission, *Symposium on Information Networks*, **III**, pp. 115-130, Polytechnic Institute of Brooklyn, April, 1954.
8. S. I. Cory, Telegraph Transmission Coefficients, *Bell Laboratories Record*, **33**, pp. 11-15, January, 1955.
9. T. A. Jones and K. W. Pfleger, Performance Characteristics of Various Carrier Telegraph Methods, *B.S.J.T.*, **25**, pp. 483-531, July, 1946.
10. C. A. Dahlbom, A. W. Horton, and D. L. Moody, Application of Multifrequency Pulsing in Switching, *Trans. A.I.E.E.*, **68**, pp. 392-396, 1949.
11. E. D. Sunde, Theoretical Fundamentals on Pulse Transmission, *B.S.T.J.*, **33**, pp. 721-788 and 987-1010, May and July, 1954.
12. I. E. Lattimer, The Use of Telephone Circuits for Picture and Facsimile Service, Long Lines Department, 1948.
13. H. Nyquist, Certain Topics in Telegraph Transmission Theory, *Trans. A.I.E.E.*, **47**, pp. 617-644, April, 1928.
14. B. M. Oliver, J. R. Pierce and C. E. Shannon, The Philosophy of PCM, *Proc. I.R.E.*, **36**, pp. 1324-1331, November, 1948.
15. R. E. Crane, J. T. Dixon, and G. H. Huber, Frequency Division Techniques for a Coaxial Cable Network, *Trans. A.I.E.E.*, **66**, pp. 1451-1459, 1947.
16. J. T. O'Leary, E. C. Blessing and J. W. Beyer, An Improved Three-Channel Carrier Telephone System, *B.S.T.J.*, **17**, pp. 162-183, January, 1938.
17. C. W. Carter, U. S. Patent No. 2,390,869.

Design, Performance and Application of the Vernier Resolver*

By G. KRONACHER

(Manuscript received May 29, 1957)

The Vernier Resolver is a precision angle transducer which, from the stand-point of performance, resembles a geared up synchro resolver, except that the step-up ratio between the mechanical angle and the electrical signal is obtained electrically.

Vernier resolvers with step up ratios of 26, 27, 32 and 33 have been designed and built.

The unit is a reluctance type, variable coupling transformer. By placing all windings on the stator, sliding contacts are eliminated. Both the stator and the rotor are laminated. Because of the averaging effect inherent in a laminated construction, the accuracy of the unit exceeds by many times the machining accuracy.

The performance of present experimental units is characterized by a repeatability of better than ± 3 seconds of shaft angle, and a standard deviation error over one full revolution of less than 10 seconds of arc.

I. INTRODUCTION

The precise measurement of an angle is a basic operation in many technical fields. The observation of stars, mapping of land, machining in the factory are all operations which require angle measurements. Of course, an angle can be measured by reading a calibrated dial. However, in automatically controlled operations the angular position of a shaft has to be sensed electrically. The instrument which performs the conversion from a mechanical angle to an electrical output is called an angle transducer. One commonly used angle transducer is the synchro resolver.

Basically this is a variable coupling transformer with one primary winding and two output windings displaced 90 degrees from each other. The variable electrical coupling is accomplished by placing the primary

* The Vernier Resolver was developed under the sponsorship of the Wright Air Development Center.

winding on the rotor and the secondary windings on the stator or vice versa. The primary winding is excited from an alternating voltage source of, say, 400 cycles per second. The amplitudes of the induced secondary voltages of the synchro resolver are ideally proportional to the sine and cosine of the rotor orientation. These two induced, amplitude modulated voltages are the resolver output.

The accuracy of commercially available synchros is, at best, three minutes of arc. Certainly, this accuracy is sufficient for many applications. In the machining of precision parts and in field applications involving the measurement of elevation and azimuth of distant targets, however, accuracies down to 10 seconds of arc are required.

One might be tempted to try to meet this requirement by merely refining the present standard synchro. However, even if this refinement were possible, it still would be a difficult task to transmit this near-perfect synchro output and also to convert it into other analog forms without losing most of the added accuracy because of noise in the system. The transmission and conversion problem can be side-stepped by going to a so-called "two speed" or "vernier" representation of the angle. This representation is obtained by using two synchros; one, the low speed synchro, is positioned directly to the particular angle and the other, the high-speed synchro, is geared up with respect to the former. The angle is now represented by two synchro outputs. Assuming perfect gears the accuracy of this system is improved by the step-up ratio in the gearing.

This approach has been adopted in the past, but unfortunately, it has major disadvantages to it. First, precision gears of better than one minute of arc are expensive, relatively large and of limited life due to wear. Second, considerable torque is required to overcome the gear friction and the inertia effect of the high speed synchro. For these reasons, it is desirable to replace the geared-up synchro by a transducer which performs the step-up between input and output electrically. The vernier resolver is such an angle transducer.

The unit is a reluctance type, variable coupling transformer. By placing all windings on the stator, sliding contacts are eliminated. Both the stator and the rotor are built up of laminations. The step-up ratio is equal to the number of teeth on the rotor lamination. Prototype units have been built with step-up ratios of 26, 27, 32 and 33. The accuracy of these units is characterized by a standard deviation error of less than 10 seconds of arc. This high degree of accuracy is due largely to the averaging effect inherent in a laminated construction. The unit may be regarded, simply, as a device which senses the average orientation of all rotor laminations with respect to the stator. Because of the great

number of laminations (one hundred in the present units) the effect of individual imperfections in laminations is greatly reduced.

In preparation for a close study of the vernier resolver we shall describe the performance of an ideal unit, and also introduce some technical terms. The output of the vernier resolver consists of two amplitude modulated voltages one of which is called the sine-voltage and the other the cosine voltage. The amplitudes of these voltages are proportional to the sine and cosine of “ n ”-times the rotor orientation. The factor “ n ” which, of course, is a function of the rotor configuration will be called the order of the resolver. We shall call the arctangent of the ratio of the secondary voltages — sine-voltage over cosine-voltage — the “signal-angle”. Furthermore, to define a positive sense of rotation and to make the signal-angle definition unambiguous, we shall assume that, with continuous positive shaft rotation, the signal-angle runs through a sequence of cycles, each going from zero to 360° . Thus, one signal-angle cycle corresponds to a shaft rotation of $(1/n)$ th; of one revolution. This angular interval is called the “vernier” interval.

II. DESIGN PRINCIPLES

2.1 *A Simplified Description*

Fig. 1 represents a simplified model of a third order vernier resolver. The unit consists of a laminated rotor with three equally spaced teeth and a laminated 4-pole-shoe stator. Each pole-shoe bears one exciting coil (not shown in the figure) and one output coil. Successive exciting coils are wound in opposite directions, connected in series, and energized from an ac source. Thus, successive pole-shoe fields alternate in phase. The two output windings each consist of two diametrical output coils connected in phase opposition.

If the rotor were a circular cylinder, the net voltage in either output winding would be zero. However, because of the three rotor teeth, the induced voltage of either output winding goes through three identical cycles per rotor revolution. Consequently, the amplitude E_c of the induced cosine-voltage e_c can be represented as a Fourier series of three times the shaft angle θ_m ,

$$E_c = E_1 \cos (3\theta_m) + E_3 \cos [3(3\theta_m)] + \cdots, \quad (1)$$

where E_1 , E_3 are the Fourier components of E_c with respect to $(3\theta_m)$.

The series is free of even harmonic terms because of the symmetry between positive and negative half-cycles. The expression for the amplitude E_s of the sine-voltage e_s is obtained by substituting $[\theta_m - (\pi/6)]$

for θ_m in (1);

$$E_s = E_1 \sin (3\theta_m) - E_3 \sin [3(3\theta_m)] + \cdots . \quad (2)$$

The magnitudes of the Fourier components depend on the design details of the unit. In a properly designed unit all higher order Fourier components are sufficiently small so that the induced signal voltages closely approximate those of an ideal third order resolver as expressed by the following equations:

$$E_c = E_1 \cos (3\theta_m), \quad (3)$$

$$E_s = E_1 \sin (3\theta_m), \quad (4)$$

$$\theta_s \equiv \tan^{-1} \frac{E_s}{E_c} = 3\theta_m, \quad (5)$$

where θ_s is the signal-angle.

2.2 Analysis of Practical Case

Fig. 2 shows an assembled unit, and typical stator and rotor laminations are illustrated in Figs. 3 and 4.

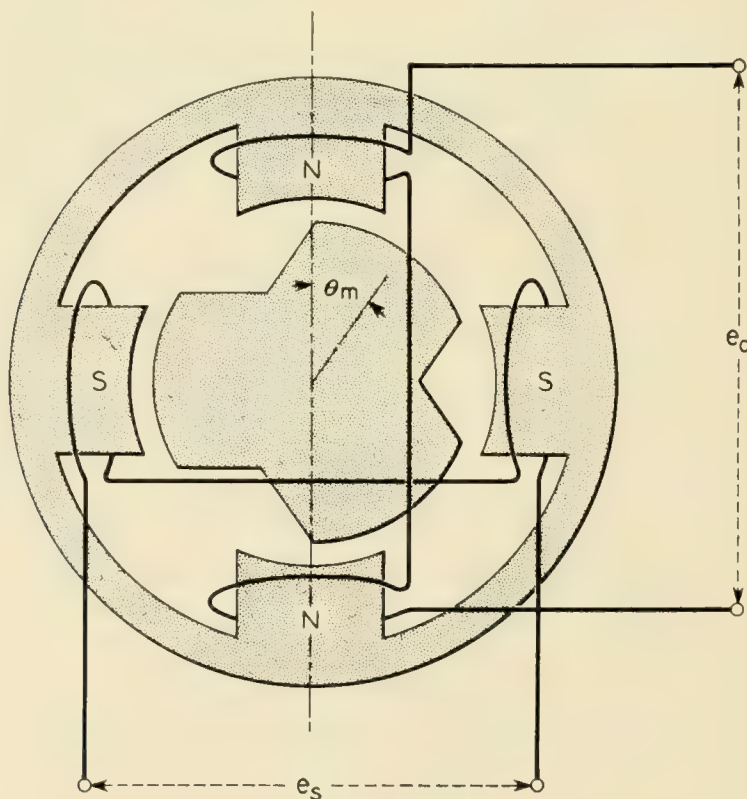


Fig. 1 — Schematic of a 3rd order vernier resolver.

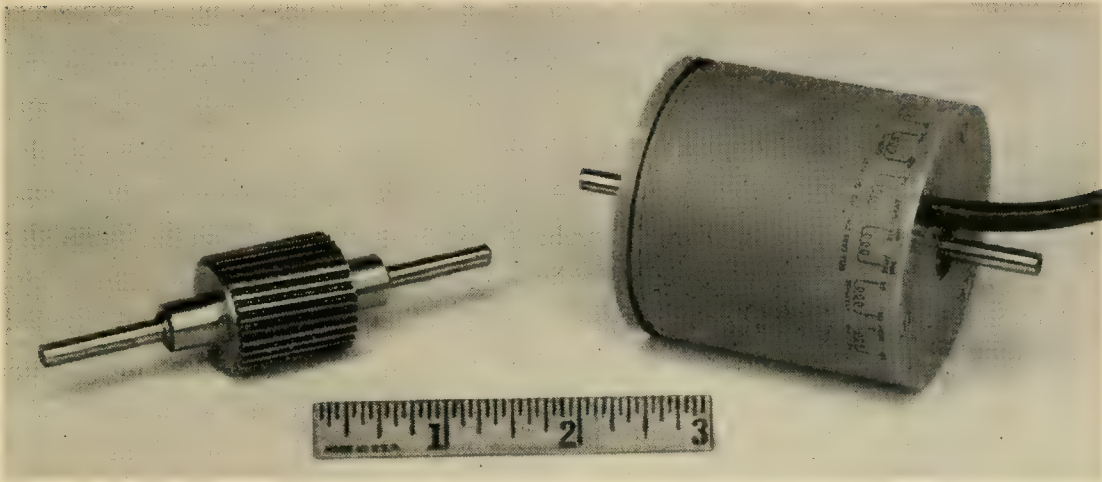


Fig. 2 — View of the assembled vernier resolver and its rotor.

The stator lamination is of ten pole-shoes, each pole-face having three teeth. All 30 teeth of the stator are equally spaced. The number of equally spaced teeth of the rotor lamination is equal to the order " n " of the resultant vernier resolver.

The exciting winding, as in the simplified model, produces ac magnetic fields alternating in phase from one pole-shoe to the next. Each of the two output windings is distributed over all ten pole-shoes.

To describe the turns distribution of these windings it is necessary to define the positive winding sense and the electrical angle of a pole-

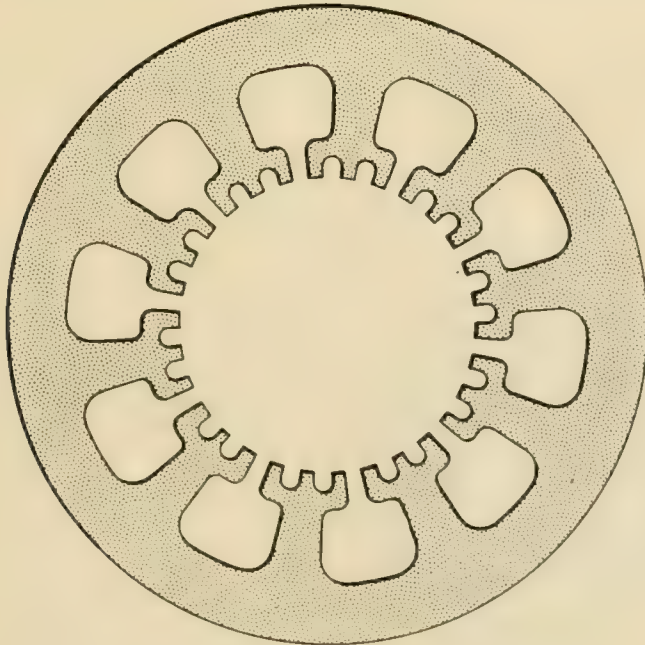


Fig. 3 — Stator lamination of the vernier resolver.

shoe. The positive winding sense for a given pole-shoe is that of the exciting coil. The electrical angle of a pole-shoe is its mechanical angle measured clockwise, with respect to a reference on the stator, multiplied by the number of rotor teeth (the order of this resolver).

The winding which produces the cosine voltage, e_c , consists of ten coils, one on each pole-shoe, connected in series. Each coil has a different number of turns depending on the pole-shoe to which it belongs. Specifically, this number of turns is equal to a design constant “ t ” multiplied by the cosine of the electrical angle of the pole-shoe. Similarly, the winding which produces the sine-voltage, e_s , consists of ten coils, one on each pole-shoe. The number of turns of each coil is equal to the same constant “ t ” multiplied by the sine of the electrical angle of the particular pole-shoe.

With α_e being the electrical angle between adjoining pole-shoes and with the electrical angle of pole-shoe No. 0 equalling zero, the turns of the coils of the cosine-winding on pole-shoes No. 0 through 9 are:

$$\begin{aligned} t_{c0} &= t \cos (0) \\ t_{c1} &= t \cos (\alpha_e) \\ &\vdots \quad \quad \vdots \\ t_{c9} &= t \cos (9\alpha_e) . \end{aligned}$$

(6)

Similarly the turns distribution, t_s , of the sine winding is:

$$\begin{aligned} t_{s0} &= t \sin (0) \\ t_{s1} &= t \sin (\alpha_e) \\ &\vdots \quad \quad \vdots \quad \quad \vdots \\ t_{s9} &= t \sin (9\alpha_e) \end{aligned}$$

(7)

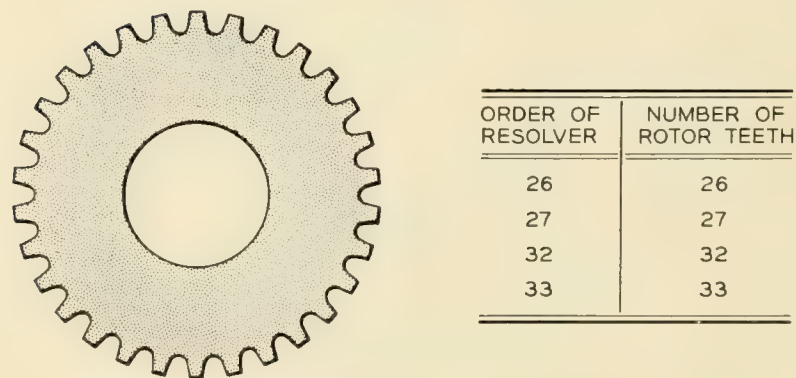


Fig. 4 — Rotor lamination of the vernier resolver.

To obtain the voltage induced in the output coils, the flux amplitude for each pole-shoe must be established. Defining the electrical angle of the rotor, θ_e , as its mechanical angle multiplied by its number of teeth and choosing θ_e to be zero when the center of a rotor tooth lines up with the center of pole-shoe No. 0, one can write for the flux amplitudes ϕ_0 through ϕ_9 of pole-shoes No. 0 through No. 9:

$$\begin{aligned} \phi_0 &= A_0 + A_1 \cos \theta_e + A_2 \cos 2\theta_e + \cdots \\ \phi_1 &= A_0 + A_1 \cos (\theta_e - \alpha_e) + \cdots \\ \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \\ \phi_9 &= A_0 + A_1 \cos (\theta_e - 9\alpha_e) + \cdots \end{aligned} \tag{8}$$

where $A_0, A_1, A_2, \dots$ are the Fourier Components of ϕ .

The amplitude, E_c , of the voltage induced in the cosine winding is the sum of the products of the pole-shoe flux, ϕ_ν , measured in [volt sec] and the coil turns, $t_{c\nu}$, multiplied by the exciting current frequency in radians per second, ω :

$$E_c = \omega \sum_{\nu=0}^9 \phi_\nu t_{c\nu} . \tag{9}$$

Substituting the values of $t_{c\nu}$ and ϕ_ν from (6) and (7) and neglecting all higher order Fourier components of ϕ , one obtains

$$E_c = \frac{10}{2} \omega A_1 \cos \theta_e . \tag{10}$$

Similarly one obtains for the amplitude E_s of the sine voltage:

$$E_s = \frac{10}{2} \omega A_1 \sin \theta_e . \tag{11}$$

As required, the two induced secondary voltages are proportional to the cosine and sine of the electrical rotor angle.

As shown in the appendix, an analysis which takes into account the higher order Fourier components of the pole-shoe flux shows that a sinusoidally distributed winding is sensitive solely to the so-called slot-harmonics. The order “ m ” of these harmonics is given by the expression

$$m = kq \pm 1 \tag{12}$$

where “ k ” is any integral positive number and q is the number of pole-shoes divided by the largest integral factor common to the number of pole-shoes and the number of rotor teeth. For instance, in the case of a

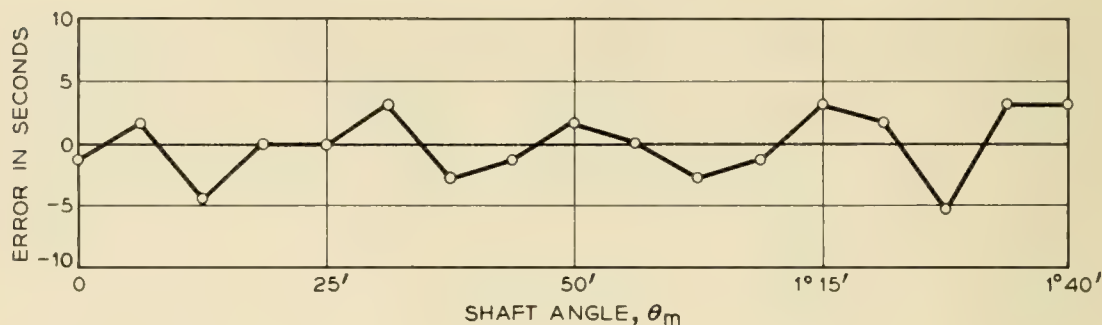


Fig. 5 — Error of a 27th order vernier resolver over $\frac{1}{8}$ vernier interval.

27th order vernier resolver this common factor is 1 and consequently the slot harmonics are of order: 9, 11, 19, 21, etc.

The effect of the slot harmonics can be reduced by the following means:

a) Selecting the dimensions as well as the number of rotor and stator teeth such as to keep the higher order flux components low.

b) Using a “skewed” rotor or stator, in which successive laminations are progressively displaced with respect to their angular orientation.

III. PERFORMANCE

Clifton Precision Products Co. built experimental resolver models of order 26, 27, 32 and 33 using the laminations shown in Figs. 3 and 4. The best results were obtained with 27th order resolvers. Their performance is described in the following sections.

3.1 Repeatability and Accuracy

The repeatability is better than ± 3 seconds of shaft angle.

Figs. 5, 6 and 7 show the error curves taken on a 27th order vernier resolver after compensating with trimming resistors for the fundamental and second harmonic error with respect to the vernier interval. (In essence, the effect of these trimming resistors is either to add or to subtract a small voltage to one or both of the resolver signals.) Fig. 8 shows an error curve before trimming.*

3.2 Temperature Sensitivity

The error introduced by a temperature change of 70°C is less than 25 seconds of shaft rotation.

* The error curves really represent the combined error of the tested resolver itself plus that of the testing apparatus.

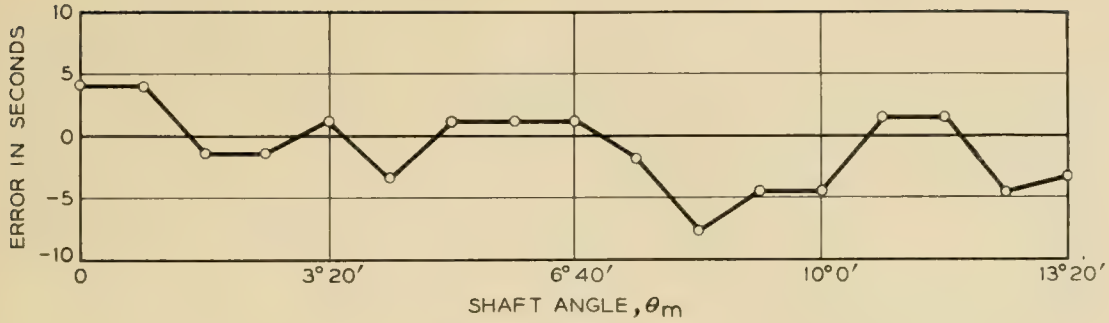


Fig. 6 — Error of a 27th order vernier resolver over one vernier interval.

It may be pointed out that the housing of the tested unit was of aluminum. A unit with a non-magnetic steel housing should be of lower temperature drift, because stator stack and housing would then have the same temperature coefficient of expansion.

3.3 Transformation Ratio, Input and Output Impedances

At maximum coupling the induced output voltage is 0.123 times the exciting voltage and is leading in phase by 6° .

The impedance of the input winding with the output windings open is $117 + j\,781$ ohms.

The impedance of the output windings with the primary winding shorted is $235 + j\,920$ ohms.

The effect of the rotor position on this impedance is hardly noticeable.

3.4 Output Signal Distortions

The harmonic content of the output signal at maximum coupling is:

fundamental 1.7 volts
 2nd harmonic 0.2 mv
 3rd harmonic 13.5 mv
 5th harmonic 5.4 mv

The harmonic content of the output signal at minimum coupling (null voltage) is:

fundamental 1.6 mv
 2nd harmonic 0.05 mv
 3rd harmonic 2.0 mv

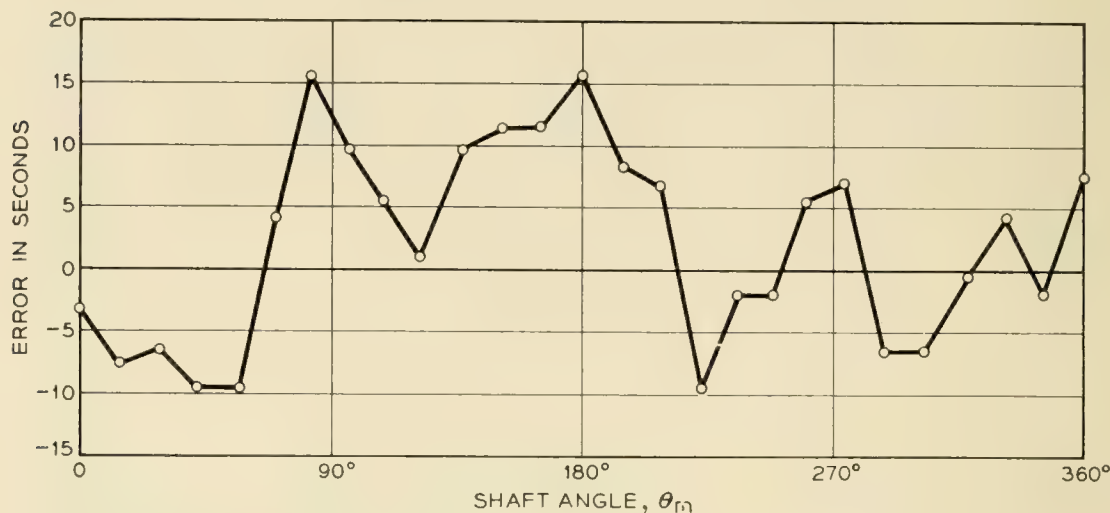


Fig. 7 — Error of a 27th order vernier resolver over one shaft revolution.

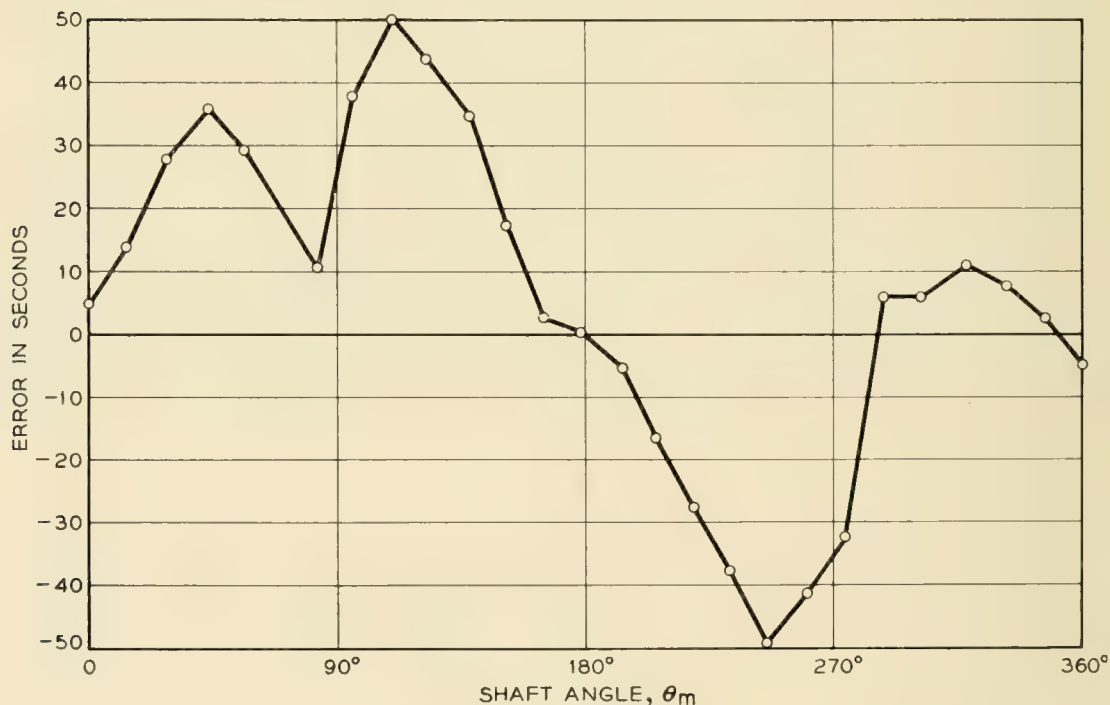


Fig. 8 — Error of a 27th order vernier resolver before trimming.

3.5 Moment of Inertia and Friction Torque

The moment of inertia of the rotor of a 27th order harmonic resolver is 63 gram cm sq.

The maximum break-away friction torque among five units was 0.027 in. oz. No change in this torque due to excitation of the unit could be detected.

IV. APPLICATION

In its application, the vernier resolver is usually directly coupled to a standard resolver or some other coarse angle transducer. Such a system which represents a variable, in this case the shaft angle, in two scales, coarse and fine, will be called a vernier system.

The following sections describe applications using the vernier resolver in an encoder, a follow-up system and an angle-reading system.

4.1 Vernier Angle Encoder

A vernier angle encoder converts a shaft angle into a pair of digital numbers, one being the coarse and the other being the vernier number. This type of encoder can be built by mechanically coupling a standard resolver directly to a vernier resolver. The outputs of the two resolvers, after encoding, represent the coarse and the vernier number.

The output of a resolver may be encoded, for instance, by the following method. The primary winding of the resolver is excited from an a-c source of, say, 400 cycles per second. The two induced secondary voltages are in phase with each other. Their amplitudes are proportional to the cosine and sine of the electrical rotor angle, θ_e .

These two amplitude modulated voltages are combined by means of two phase-shifting networks into two phase-modulated voltages. One network first advances the sine voltage by 90° and then adds it to the cosine voltage. The other network performs the same addition after retarding the sine voltage by 90° . The result is two constant amplitude voltages with relative phase shift of twice the electrical rotor angle. The time interval between the respective zero crossings of these two voltages is con-

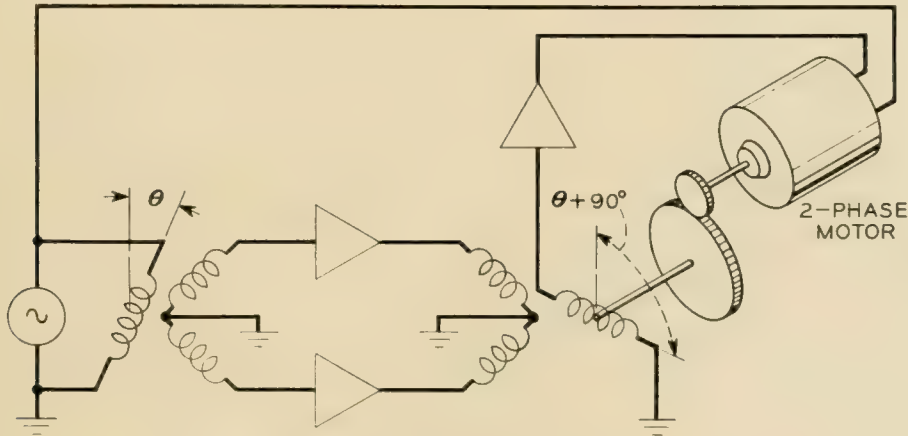


Fig. 9 — Resolver servo system.

verted into digital representation by means of an electronic stop-watch (time encoder).

4.2 *Vernier Follow-up System*

A vernier follow-up system can be built much like the present two speed synchro control-transformer system, except that the geared up synchros are replaced by vernier resolvers. Fig. 9 illustrates the vernier portion of this system. Since the output impedance of the vernier resolver is fairly large, it may be desirable to use amplifiers, as shown in Fig. 9, to energize that vernier resolver which plays the part of the control-transformer.

4.3 *Vernier Angle-Reading System*

A visual vernier angle-reading system as required to read the position of a rotary table can be built by using the output of a vernier resolver to position a standard resolver.

The coarse angle can be read as usual from calibration lines marked directly on the rotary table. The vernier angle reading is obtained by coupling a vernier resolver directly to the rotary table. The output of this resolver is used to position a standard control transformer. This control transformer will go through " n " revolutions for each revolution of the rotary table, where " n " is the order of the vernier resolver. The reading of a dial coupled either directly or through gears to the control transformer provides the vernier reading.

V. SUMMARY

Vernier resolvers of order 26, 27, 32 and 33 have been designed, built and tested.

The construction of the unit is very simple because all windings are located on the stator. The absence of brushes and slip rings makes the unit inexpensive in production and reliable in performance.

The performance of present experimental models is characterized by a repeatability of better than ± 3 seconds of arc and by a standard deviation error over one revolution of less than 10 seconds of arc.

Production units should be of even higher accuracy because better tooling fixtures would be used and minor design improvements would be incorporated.

The principal foreseeable application of the resolver lies in vernier systems. Vernier encoder, vernier servo and vernier angle-reading systems

are readily obtained by applying existing techniques to the vernier resolver.

ACKNOWLEDGEMENT

The development of the vernier resolver was undertaken under the sponsorship of the Wright Air Development Center. The work was encouraged and furthered by J. C. Lozier of Bell Telephone Laboratories. Valuable design contributions were made by J. Glass of Clifton Precision Products Co. All testing and evaluating of test results was done by T. W. Wakai of Bell Telephone Laboratories.

APPENDIX

Symbols

E	Amplitude of induced voltage
p	Number of pole-shoes
n	Number of rotor teeth
θ_e	Electrical rotor angle, equal to its geometrical angular position multiplied by n
q	p divided by largest integral factor common to n and p
n'	n divided by the same factor
ν	Pole-shoe number running from 0 to $(p - 1)$
m	Order of Fourier component representing the pole-shoe flux as a function of the electrical rotor angle, θ_e
α_e	Electrical angle between adjoining pole-shoes, equal to the geometrical angle multiplied by n
k	A number equal to zero or to any positive integer.

In accordance with equations (7) and (8) the voltage induced in the sine-winding coil on the ν th pole-shoe by the m th flux harmonic is:

$$E_{sm\nu} = \omega[A_m \cos m (\theta_e - \nu\alpha_e) t \sin (\nu\alpha_e)]. \tag{13}$$

After trigonometric transformation:

$$E_{sm\nu} = \frac{1}{2}A_m t \omega [\sin (m\theta_e - (m - 1)\nu\alpha_e) + \sin (-m\theta_e + (m + 1)\nu\alpha_e)] . \tag{14}$$

The voltage, E_{sm} , induced in the sine winding is obtained by summing the expression of (14) over all values of ν . Since $(p\alpha_e)$ is a multiple of 2π the summing of all sine terms from $\nu = 0$ to $\nu = (p - 1)$ results in zero unless the angle $(m \pm 1) \alpha_e$ is an integral multiple, k , of 2π . This condition is spelled out in the following equation:

$$(m \pm 1) \alpha_e = k2\pi. \quad (15)$$

The electrical angle α_e being the mechanical angle between successive pole-shoes divided by the number of rotor teeth is:

$$\alpha_e = \frac{2\pi n}{p}. \quad (16)$$

Dividing p and n by the largest common integral factor, one can write

$$\alpha_e = \frac{2\pi n'}{q}. \quad (17)$$

Substituting this expression into (15) and solving for m , one obtains

$$m = \frac{kq}{n'} \pm 1 \quad (18)$$

where m and k are integers or zero. Consequently, (18) is satisfied for the following values of m :

$$m = 1; \quad q \pm 1; \quad 2q \pm 1; \cdots. \quad (19)$$

The amplitude of the voltage, E_{sm} , induced in the sine winding by flux harmonics of order m , where m is specified by (19), is

$$E_{sm} = \frac{p}{2} A_m t \omega \sin (m\theta_e). \quad (20)$$

Similarly one obtains for the voltages, E_{cm} , induced in the cosine winding

$$E_{cm} = \frac{p}{2} A_m t \omega \cos (m\theta_e). \quad (21)$$

Bell System Technical Papers Not Published in this Journal

ANDERSON, O. L.,¹ CHRISTENSEN, H.,¹ and ANDREATCH, P., JR.¹

A Technique for Connecting Electrical Leads to Semiconductors, J. Appl. Phys., Letter to the Editor, **28**, p. 923, August, 1957.

ANDREATCH, P., JR., see Anderson, O. L.

BENSON, K. E., see Pfann, W. G.

Benson, K. E., see Wernick, J. H.

BIRDSALL, H. A.¹

Insulating Films, in "Digest of Literature on Dielectrics", Publication 503, National Academy of Sciences, National Research Council, **19**, pp. 209-220, June, 1957.

BOGERT, B. P.¹

Response of an Electrical Model of the Cochlear Partition with Different Potentials of Excitation, J. Acous. Soc. Am., **29**, pp. 789-792, July, 1957.

BOHM, D., see Huang, K.

BOYET, H., see Weisbaum, S.

BOZORTH, R. M.¹

The Physics of Magnetic Materials, in "The Science of Engineering Materials", John Wiley & Sons, New York, pp. 302-335, 1957.

BRADY, G. W.¹

Structure of Tellurium Oxide Glass, J. Chem. Phys., **27**, pp. 300-303, July, 1957.

¹ Bell Telephone Laboratories, Inc.

BRADY, G. W.,¹ and KRAUSE, J. T.¹

Structure in Ionic Solutions. I, J. Chem. Phys., **27**, pp. 304–308, July, 1957.

BREIDT, P., JR., see Greiner, E. S.

CHYNOWETH, A. G.,¹ and MCKAY, K. G.¹

Internal Field Emission in Silicon p-n Junctions, Phys. Rev., **106**, pp. 418–426, May 1, 1957.

CHRISTENSEN, H., see Anderson, O. L.

DACEY, G. C.,¹ and THURMOND, C. D.¹

p-n Junctions in Silicon and Germanium: Principles, Metallurgy, and Applications, Metallurgical Reviews, **2**, pp. 157–193, July, 1957.

ELLIS, W. C., see Greiner, E. S.

FRISCH, H. L.¹

The Time Lag in Nucleation, J. Chem. Phys., **27**, pp. 90–94, July, 1957.

FULLER, C. S.,¹ and REISS, H.¹

Solubility of Lithium in Silicon, J. Chem. Phys., Letter to the Editor, **27**, pp. 318–319, July, 1957.

GIBBONS, D. F.¹

Acoustic Relaxations in Ferrite Single Crystals, J. Appl. Phys., **28**, pp. 810–814, July, 1957.

GIANOLA, U. F.¹

Damage in Silicon Produced by Low Energy Ion Bombardment, J. Appl. Phys., **28**, pp. 868–873, Aug., 1957.

GITHENS, J. A.¹

The TRADIC LEPRECHAUN Computer, Proc. Eastern Joint Computer Issue, A.I.E.E. Special Publication, **T-92**, pp. 29–33, Dec. 10–12, 1956.

GLASS, M. S.¹

Straight Field Permanent Magnets of Minimum Weight for TWT — Design and Graphic Aids in Design, Proc. I.R.E., **45**, pp. 1100–1105, Aug., 1957.

¹ Bell Telephone Laboratories, Inc.

GREEN, E. I.¹

Evaluating Scientific Personnel, Elec. Engg., **76**, pp. 578-584, July, 1957.

GREINER, E. S.,¹ BREIDT, P., JR.,¹ HOBSTETTER, J. N.,¹ and ELLIS, W. C.¹

Effects of Compression and Annealing on the Structure and Electrical Properties of Germanium, J. Metals, **9**, pp. 813-818, July, 1957.

HAMMING, R. W., see Hopkins, I. L.

HARKER, K. J.¹

Nonlaminar Flow of Cylindrical Electron Beams, J. Appl. Phys., **28**, pp. 645-650, June, 1957.

HERMAN, H. C.¹

Jumbo Case Considerations, J. Patent Office Society, **39**, pp. 515-523, July, 1957.

HOBSTETTER, J. N., see Greiner, E. S.

HOPKINS, I. L.,¹ and HAMMING, R. W.¹

On Creep and Relaxation, J. Appl. Phys., **28**, pp. 906-909, Aug. 1957.

HUANG, K.,¹ BOHM, D.,⁸ and PINES, D.⁹

Role of Subsidiary Conditions in the Collective Description of Electron Interactions, Phys. Rev., **107**, pp. 71-80, July 1, 1957.

INGRAM, S. B.¹

Graduate Study in Industry — The Communications Development Training Program of the Bell Telephone Laboratories, Eng. J., **40**, pp. 993-996, July, 1957.

JACCARINO, V.,¹ SHULMAN, R. G.,¹ and STOUT, R. W.¹⁰

Nuclear Magnetic Resonance in Paramagnetic Iron Group Fluorides, Phys. Rev., Letter to the Editor, **106**, pp. 602-603, May 1, 1957.

JONES, W. D., see Turrell, G. C.

¹ Bell Telephone Laboratories, Inc.

⁸ Technion, Haifa, Israel.

⁹ Princeton University, Princeton, N. J.

¹⁰ University of Chicago, Chicago, Ill.

KIERNAN, W. J.¹

Appearance Specifications and Control Methods, Elec. Manufacturing, **60**, pp. 126-129, 294, 296, 298, July, 1957.

KRAUSE, J. T., see Brady, G. W.

LANDER, J. J., see Morrison, J.

LEGG, V. E.¹

Survey of Square Loop Magnetic Materials, Proc. Conference on Magnetic Amplifiers, A.I.E.E. Special Publication, **T-98**, pp. 69-77, Sept. 4, 1957.

LOGAN, R. A.,¹ and PETERS, A. J.¹

Diffusion of Oxygen in Silicon, J. Appl. Phys., **28**, pp. 819-820, July, 1957, Letter to the Editor.

LUKE, C. L.¹

Determination of Sulfur in Nickel by the Evolution Method, Anal. Chem., **29**, pp. 1227-1228, Aug. 1957.

MAKI, A., see Turrell, G. C.

MARCATILI, E. A.¹

Heat Loss in Grooved Metallic Surface, Proc. I.R.E., **45**, pp. 1134-1139, Aug., 1957.

MERTZ, P.¹

Information Theory Impact on Modern Communications, Elec. Engg., **76**, pp. 659-664, Aug., 1957.

MCCALL, D. W.¹

Dielectric Properties of Polythene, in "Polythene, The Technology and Uses of Ethylene Polymers", edited by A. Renfrew and P. Morgan, published for "British Plastics" by Iliffe and Sons Limited, London, 1957.

MCCALL, D. W., see Slichter, W. P.

McKAY, K. G., see Chynoweth, A. G.

¹ Bell Telephone Laboratories, Inc.

MEITZLER, A. H.¹

A Procedure for Determining the Equivalent Circuit Elements Representing Ceramic Transducers Used in Delay Lines, Proc. 1957 Electronic Components Symposium, pp. 210-219, May, 1957.

MILLER, R. C.,¹ and SMITS, F. M.¹

Diffusion of Antimony Out of Germanium and Some Properties of the Antimony-Germanium System, Phys. Rev., **107**, pp. 65-70, July 1, 1957.

MONTGOMERY, H. C.¹

Field Effect in Germanium at High Frequencies, Phys. Rev., **106**, pp. 441-445, May 1, 1957.

MORRISON, J.,¹ and LANDER, J. J.¹

The Solution of Hydrogen in Nickel Under Hydrogen Ion Bombardment, Conference Notes M.I.T. Physical Electronics Conference, pp. 102-108, June, 1957.

NIELSEN, E. G.¹

Behavior of Noise Figure in Junction Transistors, Proc. I.R.E., **45**, pp. 957-963, July, 1957.

PAYNE, R. M.⁶

Clemson Conducts School for Bell, Telephony, **152**, pp. 20-21, 50-51, June 15, 1957.

PERRY, A. D., JR.¹

Pulse-Forming Networks Approximating Equal-Ripple Flat-Top Step Response, I.R.E. Convention Record, **2**, pp. 148-153, July, 1957.

PETERS, A. J., see Logan, R. A.

PFANN, W. G.,¹ and VOGEL, F. L., JR.¹

Observations on the Dislocation Structure of Germanium Crystals, Acta Met., **5**, pp. 377-384, July, 1957.

PFANN, W. G.,¹ BENSON, K. E.,¹ and WERNICK, J. H.¹

Some Aspects of Peltier Heating at Liquid-Solid Interferences in Germanium, J. Electronics, **2**, pp. 597-608, May, 1957.

¹ Bell Telephone Laboratories, Inc.

⁶ Southern Bell Tel. & Tel. Co., Atlanta, Georgia.

PEANN, W. G.¹

Zone Melting, Metallurgical Reviews, **2**, pp. 29-76, May, 1957.

PINES, D., see Huang, K.

REISS, H., see Fuller, C. S.

RUGGLES, D. M.¹

A Miniaturized Quartz Crystal Unit for the Frequency Range 2-kc to 16-kc, Proc. 1957 Electronic Components Symposium, pp. 59-61, 1957.

SCAFF, J. H.¹

Impurities in Semiconductors, Effect of Residual Elements on the Properties of Metals, A.S.M. Special Vol., pp. 88-132, 1957.

SHULMAN, R. G., see Jaccarino, V.

SLICHTER, W. P.,¹ and McCALL, D. W.¹

Note on the Degree of Crystallinity in Polymers as Found by Nuclear Magnetic Resonance, J. Poly. Sci., Letter to the Editor, **25**, pp. 230-324, July, 1957.

SMITS, F. M., see Miller, R. C.

STEELE, A. L.⁷

The 2-5 Numbering Plan and Selection of Exchange Names, Telephony, **153**, pp. 24-25, 48, July 20, 1957.

STONE, H. A., JR.¹

Component Development for Microminiaturization, I.R.E. 1957 Convention Record, **6**, pp. 13-20, 1957.

STOUT, J. W., see Jaccarino, V.

THURMOND, C. D., see Dacey, G. C.

TURRELL, G. C.,¹ JONES, W. D.,⁴ and MAKI, A.⁴

Infrared Spectra and Force Constants of Cyanoacetylene, J. Chem. Phys., **26**, pp. 1544-1548, June, 1957.

¹ Bell Telephone Laboratories, Inc.

⁴ Oregon State College, Corvallis.

⁷ Indiana Bell Telephone Company, Indianapolis.

VAN BERGELJK, W. A.¹

The Lung Volume of Amphibian Tadpoles, *Science*, **126**, p. 120, July 19, 1957.

VOGEL, F. L., JR., see Pfann, W. G.

WEINREICH, G.¹

Ultrasonic Attenuation by Free Carriers in Germanium, *Phys. Rev.*, Letter to the Editor, **107**, pp. 317-318, July 1957.

WEINREICH, G.,¹ and WHITE, H. G.¹

Observation of the Acoustoelectric Effect, *Phys. Rev.*, Letter to the Editor, **106**, pp. 1104-1106, June 1, 1957.

WEISBAUM, S.,¹ and BOYET, H.¹

Field Displacement Isolators at 4-, 6-, 11- and 24-Kmc, *Trans. I.R.E.-PGMTT*, **MTT-5**, pp. 194-198, July, 1957.

WEISS, M. T.¹

Quantum Derivation of Energy Relations Analogous to Those for Nonlinear Reactances, *Proc. I.R.E.*, Letter to the Editor, **45**, pp. 1012-1013, July, 1957.

WEISS, M. T.¹

A Solid State Microwave Amplifier and Oscillator Using Ferrites, *Phys. Rev.*, Letter to the Editor, **107**, pg. 317, July 1, 1957.

WERNICK, J. H.,¹ and BENSON, K. E.¹

Zone Refining of Bismuth, *J. Metals*, **9**, p. 996, July, 1957.

WERNICK, J. H., see Pfann, W. G.

WHITE, H. G., see Weinreich, G.

¹ Bell Telephone Laboratories, Inc.

Recent Monographs of Bell System Technical Papers Not Published in This Journal*

AARON, M. R.

Use of Least Squares in System Design, Monograph 2828.

ANDREATCH, P., JR. and THURSTON, R. N.

Disk-Loaded Torsional Wave Delay Line, Monograph 2827.

BOND, W. L., see McSkimin, H. J.

BOOTHBY, O. L., see Williams, H. J.

BOYET, H. see Seidel, H.

BOYET, H. and SEIDEL, H.

Analysis of Nonreciprocal Effects in N-wire, Ferrite-Loaded Transmission Line, Monograph 2829.

BUCK, T. M. and McKIM, F. S.

Depth of Surface Damage Due to Abrasion on Germanium, Monograph 2805.

BURKE, P. J.

Output of a Queuing System, Monograph 2766.

BURNS, F. P.

Piezoresistive Semiconductor Microphone, Monograph 2830.

CIOFFI, P. P.

Rectilinearity of Electron-Beam Focusing Fields from Transverse Determinations, Monograph 2844.

DAVID, E. E., JR.

Signal Theory in Speech Transmission, Monograph 2831.

* Copies of these monographs may be obtained on request to the Publication Department, Bell Telephone Laboratories, Inc., 463 West Street, New York 14, N. Y. The numbers of the monographs should be given in all requests.

DEWALD, J. F.

Transient Effects in Ionic Conductance of Anodic-Oxide Films,
Monograph 2767.

EASLEY, J. W.

Effect of Collector Capacity on Transient Response of Junction Transistors, Monograph 2832.

GREEN, E. I.

Evaluating Scientific Personnel, Monograph 2846.

HERRMANN, G. F.

Transverse Scaling of Electron Beams, Monograph 2839.

KARP, A.

Backward-Wave Oscillator Experiments at 100 to 200 Kilomegacycles,
Monograph 2833.

LAW, J. T. and MEIGS, P. S.

Rates of Oxidation of Germanium, Monograph 2834.

LUMSDEN, G. Q.

Wood Poles for Communication Lines, Monograph 2818.

MARRISON, W. A.

A Wind-Operated Electric Power Supply, Monograph 2837.

McKIM, F. S., see Buck, T. M.

McSKIMIN, H. J. and BOND, W. L.

Elastic Moduli of Diamond, Monograph 2793.

MEIGS, P. S., see Law, J. T.

READ, M. H., see Van Uitert, L. G.

SCHNETTLER, F. J., see Van Uitert, L. G.

SEIDEL, H.

Ferrite Slabs in Transverse Electric Mode Waveguide, Monograph 2798.

SEIDEL, H., see Boyet, H.

SEIDEL, H. and BOYET, H.

Polder Tensor for Single-Crystal Ferrite with Small Cubic Symmetry Anisotropy, Monograph 2843.

SHERWOOD, R. C., see Williams, H. J.

SHOCKLEY, W.

Transistor Physics, Monograph 2836.

SLICHTER, W. P.

Nuclear Magnetic Resonance in Some Fluorine Derivatives of Polyethylene, Monograph 2819.

SWANEKAMP, F. W., see Van Uitert, L. G.

THURSTON, R. N., see Andreatch, P., JR.

VAN UITERT, L. G.

Magnesium-Copper-Manganese-Aluminum Ferrites for Microwave Uses, Monograph 2799.

VAN UITERT, L. G.

Magnetic Induction and Coercive Force Data on Members of Series $\text{BaAl}_x\text{Fe}_{12-x}\text{O}_{19}$, Monograph 2801.

VAN UITERT, L. G.

Effects of Annealing on Saturation Induction of Ferrites With Nickel or Copper, Monograph 2845.

VAN UITERT, L. G., READ, M. H., SCHNETTLER, F. J., and SWANEKAMP, F. W.

Permanent Magnet Oxides Containing Divalent Metal Ions, Monograph 2841.

WALKER, L. R.

Magnetostatic Modes in Ferromagnetic Resonance, Monograph 2800,

WERTHEIM, G. K.

Energy Levels in Electron-Bombarded Silicon, Monograph 2840.

WILLIAMS, H. J., SHERWOOD, R. C., and BOOTHBY, O. L.

Magnetostriction and Magnetic Anisotropy of MnBi, Monograph 2838.

WILLIS, F. H.

Some Results with Frequency Diversity in a Microwave Radio System, Monograph 2842.

Contributors to This Issue

ANDREW H. BOBECK, B.S.E.E., 1948; M.S.E.E., 1949, Purdue University; Bell Telephone Laboratories, 1949-. Since completing the Laboratories' Communications Development Training Program in 1952 Mr. Bobeck has been engaged in the design of both communications and pulse transformers and more recently in the design of solid state memory devices. Member I.R.E., Eta Kappa Nu and Tau Beta Pi.

B. C. BELLOW, JR., B.S., Cornell University, 1936; General Electric Co., 1936-39; Bell Telephone Laboratories, 1939-. From 1939 to 1941 Mr. Bellows was engaged in engineering trial installations of telephone equipment, particularly multi-channel coaxial cable equipment. During World War II he specialized in the mechanical design and engineering of airborne radars. From 1945 to 1957 he was engaged in the design of circuits and equipment for point-to-point microwave radio relay systems for telephone and television transmission. On May 1, 1957 he was named Transmission Measurement Engineer. Member Eta Kappa Nu and Phi Kappa Phi.

CHARLES A. DESOER, Dipl. Ing., University of Liege (Belgium), 1949; Sc.D. Massachusetts Institute of Technology, 1953; Bell Telephone Laboratories, 1953-. Since joining the Laboratories Mr. Desoer has specialized in linear and transistor network development in the Transmission Networks Development Department. Senior Member I.R.E.

R. SHIELDS GRAHAM, B.S., University of Pennsylvania, 1937; Bell Telephone Laboratories, 1937-. His work has been with the design of equalizers, electrical wave filters and similar apparatus for use on long-distance coaxial cable circuits, microwave systems and both telephone and television transmission. During World War II, Mr. Graham designed circuits for electronic fire control computers for military use. He also developed methods for computing network and similar problems on a digital relay computer. He presently supervises the video and intermediate frequency network group. He is a senior member of the I.R.E., and a member of Tau Beta Pi, and Pi Mu Epsilon.

GERALD KRONACHER, Dipl. Eng., Federal Institute of Technology, Zurich, Switzerland, 1937; Assistant Professor, Federal Institute of Technology, 1938; mining engineer, Bolivia, 1939–1946; General Electric Company, 1946–1948; Air Associates, Inc., 1948–1951; Arma Corporation, 1951–1953; Bell Telephone Laboratories, 1953–. Since joining the Laboratories Mr. Kronacher has been associated with the Military Systems Engineering Department studying input and output problems for digital computers. He is the author of many published technical articles.

PIERRE MERTZ, A. B., 1918; Ph.D., 1926, Cornell University; American Telephone and Telegraph Company, 1919–1921, 1926–1934; Bell Telephone Laboratories, 1935–. Mr. Mertz's work with the Bell System has been concerned primarily with transmission problems relating to telephotography and television. Since 1950 Mr. Mertz has acted as a consultant in the Systems Engineering Department on such projects as micro-image readers and commercial and military data transmission problems. Fellow of the I.R.E. and the Society of Motion Picture and Television Engineers; member, American Physical Society, Optical Society of America and the Inter-Society Color Council.

DOREN MITCHELL, B.S., Princeton University, 1925; American Telephone and Telegraph Company, 1925–1934; Bell Telephone Laboratories, 1934–. Mr. Mitchell's early work with the Bell System was concerned with field studies of transmission on long telephone circuits and radio circuits, including supervision of the initial operation of the New York to Buenos Aires radio-telephone circuit. Until 1942 Mr. Mitchell worked on voice operated devices of various kinds including compandors, echo suppressors and automatic switching devices. During World War II he participated in military projects involving transmission systems and problems of laying wire from airplanes. Since the war Mr. Mitchell has been primarily concerned with radio systems. In 1955 he was appointed a Special Systems Engineer supervising a data transmission system for the SAGE project, and planning other special services involving radio. Mr. Mitchell has been granted over seventy patents. Member I.R.E.

R. C. PRIM, B.S. in E.E., University of Texas, 1941; A.M., Ph.D., Princeton University, 1949; General Electric Company, 1941–1944; Naval Ordnance Laboratory, 1944–1948; Bell Telephone Laboratories, 1949–. Since joining the Laboratories Mr. Prim has been a member of the Mathematical Research Department engaged in research and con-

sultation in the fields of theoretical mechanics, solid state electronics, aerial warfare and activities analysis. In 1955 he was placed in charge of a sub-department concerned with Computing and Theoretical Mechanics, and is presently in charge of the Communications Fundamentals sub-department. Member American Mathematical Society, American Physical Society, Tau Beta Pi and Sigma Xi.

WERNER ULRICH, B.S., 1952; M.S., 1953; Eng. Sc.D., 1957, Columbia School of Engineering; Bell Telephone Laboratories, 1953-. Mr. Ulrich's first assignment was on the design of an input circuit for an electronic memory and control device. Subsequently he was engaged in the design of logical circuits for electronic controls. Since 1954, he has been working on automatic testing and maintenance facilities for electronic switching systems. Mr. Ulrich is a member of the I.R.E. and Tau Beta Pi.

Index to Volume XXXVI

A

AF *See* United States Air Force

Activation of Electrical Contacts by Organic Vapors (L. H. Germer, J. L. Smith) 769-812

Agamemnon (cable ship) 303

AIR FORCE *See* United States Air Force

ALASKAN TELEPHONE CABLE 168

Alignment *See* Misalignment

ALLENTOWN PLANT (WESTERN ELECTRIC) 107, 123-25

ALTERNATOR SET
two-motor carrier, L-type 140

AMERICAN TELEPHONE AND TELEGRAPH COMPANY 3, 14-15

AMPLITUDE MODULATION
data transmission systems 1451-86

AMSTERDAM, HOLLAND 9

ANALYSIS
combinatorial, error correcting coding 517-35

Angell, Miss D. T. 1033

ANGLE(S)
measurement
vernier resolver 1487-1500

ANGLO-IRISH CABLE 179

Anson, H. W. 1093

ANTENNA
radio transmission
beyond the horizon 639-40
diagrams 599
height 597;
transmission loss 593-97

ARMOR, repeaters, transatlantic telephone cable 58

ARRAY *See* Memory Arrays

ATLANTIC CABLE 1-2, 4, 303

ATLANTIC OCEAN
Mid-Atlantic Ridge 1066-68
North Atlantic, *see* North Atlantic Ocean
physiographic diagram inside rear cover Sept.

ATMOSPHERE, dielectric constant 603, 627

ATTENUATION
Newfoundland-Nova Scotia link 221-23
nonlinear, FM signal 879-89
waveguide coupling 392

Aulock, Wilhelm von
biographical material 591
Measurement of Dielectric and Magnetic Properties of Ferromagnetic Materials at Microwave Frequencies 427-48

AZORES map 8, 294, 296

B

BOD TEST *See* Test: biochemical oxygen demand

BSTJ *See* Bell System Technical Journal

BACTERIA 1097-1127

Baldwin, J. A. 1337

Bampton, J. F.
biographical material 338
System Design for the Newfoundland-Nova Scotia Link 217-44

BANDWIDTH
transatlantic telephone cable
North Atlantic link 34-35

BATTERY *See* Storage Battery

BEAM *See* Electron Beam

Bechofer, R. E. 576

BELL SYSTEM TECHNICAL JOURNAL
advisory board, *see* inside front covers
Bullock, W. D., editor 710
editorial committee, *see* inside front covers
editorial staff, *see* inside front covers

BELL TELEPHONE LABORATORIES 4, 57-58, 163

Bellows, B. C., Jr.
biographical material 1511
Experimental Transversal Equalizer for TD-2 Radio Relay System 1429-50

- Beneš, Vaclav E.
 biographical material 1045
Fluctuations of Telephone Traffic 965-73
Sufficient Set of Statistics for a Telephone Exchange Model 939-64
- BERNE, SWITZERLAND 9
- Binary Block Coding* (S. P. Lloyd) 517-35
- BINARY DIGIT *See* Bit
- BINOMIAL PROCESSES 537-76
- BIOCHEMICAL OXYGEN DEMAND TEST
See Test
- Biskeborn, M. C.
 biographical material 338
Cable Design and Manufacture for the Transatlantic Submarine Cable System 189-216, 496
- BIT (BINARY DIGIT)
 AM leased-line transmission 1451-86
 twistor devices 1336
- Bleicher, E. 576
- BLOCK CODING *See* Code
- Bobeck, Andrew H.
 biographical material 1511
New Storage Element Suitable for Large Sized Memory Arryas—the Twistor 1319-40
- Bobis, S. 1449
- BORER(s), marine 194
- BOSTON map 8
- Boyd, Richard C.
 biographical material 588
New Carrier System for Rural Service 349-90
- Boyet, H. 426
- Braga, F. J.
 biographical material 338
Repeater Design for the North Atlantic Link 69-101
- Bridgers, H. E. 1004
- BRITISH POST OFFICE, submarine cables, 3-5, 14-15, 57-58, 245
- Brockbank, R. A.
 biographical material 339
Repeater Design for the Newfoundland-Nova Scotia Link 245-76
- BRUSSELS 9
- Buckley, O. E. 67
- BUILDINGS, radio transmission and 613-14
- Bullington, Kenneth
 biographical material 828
Radio Propagation Fundamentals 593-626
- Bulloch, W. D., B.S.T.J. editor 710
- Burke, P. J. 964
- C**
- C.C.I.F. *See* International Consultative Committee on Telephony
- CABLE(s)
 Alaska telephone, *see* Alaskan Telephone Cable
 Anglo-Irish, *see* Anglo-Irish Cable
 Hawaii telephone, *see* Hawaiian Telephone Cable
 short-circuits, Poisson patterns 1005-33
 submarine, *see* Submarine Cable
 transatlantic telegraph (1866), *see* Atlantic Cable
 transatlantic telephone, *see* Transatlantic Telephone Cable
 trunks, *see* Trunk
Cable Design and Manufacture for the Transatlantic Submarine Cable System (M. C. Biskeborn, H. C. Fischer, A. W. Lebert) 189-216, 496
- CABLE LAYING 13
 dynamics 1129-1207
 kinematics 1129-1207
 laying effect 43-44
 methods, early 303
 oceanography 1049
 strains 13-14
 transatlantic telephone cable 293-326
- CABOT STRAIT 3
- CANADIAN COMSTOCK CO., LTD. 244
- CANADIAN OVERSEAS TELECOMMUNICATION CORPORATION 3, 7, 244
- CAPACITANCE
 geometries, pressure coefficients 485-95
 submarine cables 485
- CAPACITOR
 carrier, P1 367-68
 mica, repeater, flexible, North Atlantic link 125-26
 parallel plate
 capacitance, pressure coefficients 485-95

CARBON

- activating
- contacts, electrical 769-812

CARRIER(s)

- history 350

L-type

- alternator set 140

P1 349-90

- capacitors 367-68
- channels 357-59
- compandors 357
- dialing 361-62
- equipment arrangements 369-76
- filters 365-66
- inductors 366-67
- installation 380-90
- maintenance 375
- miniaturization 370
- networks 369
- parameters 351-65
- power supply 376-80
- printed circuitry 371
- repeaters 362-65, 373-75
- ringing 359-61
- signaling 359
- terminals
 - block diagram 354
 - mounting 373-75
- testing 375
- transformers 368-69
- transistors 349-90
- transmission plan 351-55
- trunks 350

CASTING RESINS

- BOD test 1099
- marine conditions 1095-1127

CATALINA ISLAND 13

CENTRAL OFFICE, service range 350

Chaffee, J. G. 1449

CHANNEL(s)

- carrier, P1 357-59
- human, information rate, reading rates and 497-516
- noisy, error correction codes 1341-88
- North Atlantic link 38

Character of Waveguide Modes in Gyromagnetic Media (H. Seidel) 409-26*Circular Electric Wave Transmission in a Dielectric-Coated Waveguide* (H.-G. Unger) 1253-78*Circular Electric Wave Transmission through Serpentine Bends* (H.-G. Unger) 1279-92CIRCULAR WAVEGUIDE *See* WaveguideCIRCUIT *See* subhead circuit under names of specific equipment and apparatus, e.g. Repeater: flexible: North Atlantic link: circuit; *Also see* Printed Circuitry; Short Circuit

CLAPP, WILLIAM F., LABORATORIES 1115-21

CLARENVILLE, NEWFOUNDLAND 2, 9, 29, 49, 57, 140, 145, 147, 150, 164-65, 217, 219, 221, 246, 248, 293, 301, 317-18, 323; *map* 8, 218, 300CLARENVILLE-OBAN LINK *See* North Atlantic LinkCLARENVILLE-SYDNEY MINES LINK *See* Newfoundland-Nova Scotia Link

CLIFTON PRECISION PRODUCTS COMPANY 1494, 1499

CODE(s), CODING

- block, binary 517-35
- error correcting 517-35
 - binary, non 1341-88
 - geometric concept 1343-44
 - non-binary 1341-88
 - purpose 1341-43
- Reed-Muller 1341

Coincidences in Poisson Patterns (E. N. Gilbert, H. O. Pollak) 1005-33*Cold Cathode Gas Tubes for Telephone Switching Systems* (M. A. Townsend) 755-68

COMBINATORIAL ANALYSIS

- error correcting coding 517-35

COMMUNICATION

- channel, *see* Channel

CONNECTION, shortest, network 1389-1401

COMPANDING

- improvement 671, 688-90
- instantaneous
 - signals, quantized 653-709
- carrier, P1 357

COMPANY, TELEPHONE *See* Operating Companies

CONDUIT

- metal, testing 737-54
- round, strength requirements 737-54

CONDUIT (*Cont.*)

underground

clay, vitrified 737

loads 737

metal, testing 737-54

round

strength requirements 742-54

CONSOL (navigation system) 1050

CONTACT(S)

relay

arcing 769-812

erosion, by vapors 769-812

CONTINENTAL SHELF, North Atlantic Ocean 1063-64

Cook, Madeline L. 1126

COOPERATION *See* International Cooperation

COPENHAGEN 9

COPPER TUBING

repeater, flexible, North Atlantic link 114-15

CORONA

power supply, transatlantic telephone cable 159

COUPLED-WAVE TRANSDUCER *See* Transducer

COUPLER, waveguide

attenuation 392

design 394-401

Crawford, Arthur B.

biographical material 828

Reflection Theory for Propagation beyond the Horizon 627-44

CROSSTALK, transatlantic telephone cable 161, 230, 243

CRYSTAL

ferrites, *see* Ferrite

quartz

repeater, flexible, North Atlantic link 120-23

CUBA 13

Curtis, Harold E. 889

biographical material 828

Interchannel Interference due to Klystron Pulling 645-52

D

DC *See* Direct Current

Dagnall, C. H. 1449

DATA TRANSMISSION

AM systems 1451-86

error correcting codes, non-binary 1341-88

leased-line services, transmission aspects 1451-86

mutilation 1342

Dawson, Robert W.

biographical material 588

Experimental Dual Polarization Antenna Feed for Three Radio Relay Bands 391-408

DAYTONA BEACH, FLORIDA, test site 1115-21

DECCA NAVIGATOR 1050

DECODER, error correcting codes, non-binary 1341-88

DeCoste, J. B. 1126

DEFENSE WORK *See* Military Communications

De Hoff, Barbara 448

Depew, C. 768

Design, Performance and Application of the Vernier Resolver (G. Kronacher) 1487-1500

Desoer, Charles A. 156-58

biographical material 1511

Network Containing a Periodically Operated Switch Solved by Successive Approximations 1403-28*Determination of Pressure Coefficients of Capacitance for Certain Geometries* (D. W. McCall) 485-95

DIAL TELEPHONE, DIALING

cable, *see* Cablecarrier, *see* Carrierchannels, *see* Channel

circuit

data transmission services 1451-86

companies, *see* Operating Companiesexchange, *see* Telephone Exchangeleased-lines, *see* Leased-Line Servicesrepeaters, *see* Repeaterrural, *see* Rural Telephone Servicetelephone exchange, *see* Telephone Exchangetraffic, *see* Traffictransmission, *see* Transmissiontrunks, *see* Trunk

- DIELECTRIC, inhomogeneous, waveguide, circular, curved 1209-51
 DIELECTRIC-COATED WAVEGUIDE *See* Waveguide
 DIELECTRIC CONSTANT
 capacitance, pressure coefficients 485
 DIGIT, binary *See* Bit
 DIRECT CURRENT
 transatlantic telephone cable 139-62
 DISTORTION
 FM signal, noise modulated 879-89
 TD-2 radio relay system 1429-50
Distortion Produced in a Noise Modulated FM Signal by Non-Linear Attenuation and Phase Shift (S. O. Rice) 879-98
 Doba, S. 889
Dynamics and Kinematics of the Laying and Recovery of Submarine Cable (E. E. Zajac) 1129-1207
- E**
- E.P.I. (electronic position indicator) 1050
 EARTHQUAKES, ocean bottom 1086-88
 EASTERN TELEPHONE AND TELEGRAPH COMPANY 7, 244
 Ebbe Grace L. 1449
 ECHO, transatlantic telephone cable 21
 ECHO SOUNDING 1051-52, 1055
 EFFICIENCY
 electron tubes, transatlantic telephone cable 3
 ELASTOMER(S)
 BOD test 1099
 marine conditions 1905-1127
 ELECTRIC WAVE *See* Wave
 ELECTRICAL ATTENUATION *See* Attenuation
 ELECTRICAL CAPACITANCE *See* Capacitance
 ELECTRICAL CAPACITOR *See* Capacitor
 ELECTRICAL CONTACTS *See* Contact
 ELECTRICAL DISTORTION *See* Distortion
 ELECTRICAL FILTERS *See* Filter
 ELECTRICAL INDUCTOR *See* Inductor
 ELECTRICAL INTERFERENCE *See* Interference
 ELECTRICAL LOADING *See* Loading
 ELECTRICAL LOSS *See* Net Loss; Transmission Loss
 ELECTRICAL NETWORK *See* Network
 ELECTRICAL NOISE *See* Noise
 ELECTRICAL TRANSFORMER *See* Transformer
Electrically Operated Hydraulic Control Valve (J. W. Schaefer) 711-36
 ELECTRON BEAM
 (in) magnetic field, longitudinal
 noise spectrum 831-78
 growing noise 831-53
 UHF 855-78
 noise, electron beam, relation 832
 Pierce-type, noise 833
 ELECTRON TUBE
 6P10 179-80; *illus* 165
 6P12 180-88
 electrical characteristics 184-86
 lifetime 185-86
 tests 186-88
 175HQ 163-79
 cathode assembly *illus* 169
 electrical characteristics 171-77
 fabrication 177-78
 heater *illus* 169
 mechanical features 168-78
 reliability 178
 selection 177-78
 SP61 179-80
 gas diode
 by-pass 88-90; *illus* 89
 gas discharge
 characteristics 755-65
 switching systems 755-68
 submarine cables, performance requirements 163-88
 traveling-wave, noise 831
 See also Klystron Pulling
Electron Tubes for the Transatlantic Cable System (M. F. Holmes, J. O. McNally, G. H. Metson, E. A. Veazie) 163-88
 ELECTRONIC POSITION INDICATOR 1050
 Elmendorf, C. H.
 biographical material 1317
 Oceanographic Information for Engineering Submarine Cable Systems 1047-93

- Emling, J. W.
 biographical material 339
 *Transatlantic Telephone Cable System—
 Planning and Over-All Performance*
 7-27
- ENCODER, error correcting codes, non-
 binary 1341-88
- ENGINEERING, traffic 940
- ENGLAND map 294, 296
- EQUALIZER(S)
 repeater, flexible
 North Atlantic link 99-100
 shore
 transatlantic telephone cable 46-49
 transverse
 TD-2 radio relay system 1429-50
- EQUIPMENT MINIATURIZATION *See* Mini-
 aturization
- ERROR CORRECTING CODES *See* Code
- Ewing, Maurice 1093
- Experimental Dual Polarization Antenna
 Feed for Three Radio Relay Bands*
 (R. W. Dawson) 391-408
- Experimental Transversal Equalizer for
 TD-2 Radio Relay System* (B. C.
 Bellows, Jr., R. S. Graham) 1429-50

F

- 4B TRANSISTOR *See* Transistor
- 5B TRANSISTOR *See* Transistor
- FM *See* Frequency Modulation
- FADING, radio transmission 600-01
- Fedukowicz, W. 1093
- Feher, George
 biographical material 588
 *Sensitivity Considerations in Micro-
 wave Paramagnetic Resonance Ab-
 sorption Techniques* 449-84
- FERRITE
 microwave region, dielectric proper-
 ties, measurement 427-48
 parameters 428-30
- FERRITE LOADED WAVEGUIDE *See* Wave-
 guide
- FERROMAGNETIC MATERIALS *See* Ferrite
- Field, Cyrus 293
- FIELD *See* Magnetic Field
- FILTER(S), carrier, P1 365-66

- Fischer, H. C.
 biographical material 339
 *Cable Design and Manufacture for the
 Transatlantic Submarine Cable Sys-
 tem* 189-216, 496
- Fletcher, R. C. 426, 483, 1337
 Fluctuations of Telephone Traffic
 (V. E. Beneš) 965-73
- Fog, and radio transmission 602
- Foster, H. 1093
- FRANCE map 294, 296
- FRANKFURT, Germany 9
- Fraser, John M.
 biographical material 340
 *System Design for the North Atlantic
 Link* 29-68
- FREQUENCY, transatlantic telephone
 cable 18, 24-26
- FREQUENCY MODULATION
 interference, interchannel
 klystron pulling 645-52
- FRESNEL ZONES 597-99
- Friis, Harald T.
 biographical material 828
 *Reflection Theory for Propagation be-
 yond the Horizon* 627-44

G

- GAS DIODE TUBE *See* Electron Tube
- GAS DISCHARGE TUBE *See* Electron
 Tube
- GEIGER COUNTERS, Poisson patterns
 1005-33
- GENERALIZED TELEGRAPHIST'S EQUA-
 TIONS
 waveguide, circular, curved
 dielectric, inhomogeneous 1209-51
- Gere, E. 483
- Germer, Lester H.
 *Activation of Electrical Contacts by
 Organic Vapors* 769-812
 biographical material 829
- Geschwind, S. 483
- Gibson, W. C. 1126
- Gilbert, E. N. 964
 biographical material 1045
 Coincidences in Poisson Patterns
 1005-33
- Gilbert, J. J. 67

- Glaser, J. L. 698
 GLASGOW, Scotland 9
 Gleichmann, T. F.
 biographical material 340
 Repeater Design for the North Atlantic Link 69-101
 Graham, R. Sheils
 biographical material 1511
 Experimental Transversal Equalizer for TD-2 Radio Relay System 1429-50
Great Eastern (cable ship) 293, 303
 GREENLAND map 8
 Griffith, R. G.
 biographical material 340
 Transatlantic Telephone Cable System—Planning and Over-All Performance 7-27
 Grismore, O. D. 495
 GUIDE See Waveguide
 GUIDED MISSILES See Nike
 Gumley, R. H. 812
 GUN See Electron Gun
 Gupta, S. S. 576

H

- H.M.T.S. Monarch* See *Monarch*
 Hagelbarger, D. W. 1033
 Hale, A. L. 1093
 Halsey, R. J.
 biographical material 341
 System Design for the Newfoundland-Nova Scotia Link 217-44
 Transatlantic Telephone Cable System—Planning and Over-All Performance 7-27
 Hamming, R. W. 535
 HAWAIIAN TELEPHONE CABLE 168
 HAWTHORNE WORKS (WESTERN ELECTRIC) 107
 Heezen, Bruce C.
 biographical material 1317
 Oceanographic Information for Engineering Submarine Cable Systems 1047-93
 Heffner, William W.
 biographical material 341
 Repeater Production for the North Atlantic Link 103-38
 Heskett, H. E. 405

- High-Voltage Conductivity-Modulated Silicon Rectifier* (M. B. Prince, H. S. Veloric) 975-1004
 HILLSIDE PLANT (WESTERN ELECTRIC) 103-38
 Hipkins, Renee 512
 Hogg, David C.
 biographical material 829
 Reflection Theory for Propagation beyond the Horizon 627-44
 Holdaway, V. L. 768
 Holmes, M. F.
 biographical material 341
 Electron Tubes for the Transatlantic Cable System 163-88
 Hoth, D. F. 698
 Howard, John D.
 biographical material 588
 New Carrier System for Rural Service 349-90
 Huyett, Marilyn J.
 biographical material 589
 Selecting the Best One of Several Binomial Populations 537-76
 HUMAN CHANNEL See Channel

I

- ICELAND map 8
 INDUCTOR
 carrier, P1 366-67
 repeater, flexible
 North Atlantic link 127-29
 INFORMATION RATE
 channel, human 497-516
 prose 501-04
 reading rates 497-516
 word length 500
 vocabulary size 409
 vocoder 497
 INFORMATION STORAGE
 twistor 1319-40
 INHOMOGENEOUS DIELECTRIC See Dielectric
 INSPECTION
 repeater, flexible, transatlantic telephone cable 131-38
 INSTALLATION
 carrier, P1 380-90

Instantaneous Companding of Quantized Signals (B. Smith) 653-709

INTEGRITY (components) 31, 33

Interchannel Interference due to Klystron Pulling (H. E. Curtis, S. O. Rice) 645-52

interchannel, frequency modulation

klystron pulling 645-52

power spectrum 647-48

transatlantic telephone cable power supply 161

INTERNATIONAL CONSULTATIVE COMMITTEE ON TELEPHONY standards 17

INTERNATIONAL COOPERATION 7-8, 14-15, 27, 246, 326

IONOSPHERE, radio transmission 618-23

IOWA ENGINEERING EXPERIMENT STATION

conduit, underground 737-54

IRELAND map 294, 296

J

J-7 TRANSDUCER *See* Transducer: electrohydraulic

Jack, John S.

biographical material 341

Route Selection and Cable Laying for the Transatlantic Cable System 293-326

Jacobs, O. B. 67

Jensen, R. A. 1337

Jervey, W. T. 754

JUTE, in BOD test 1099

K

Kankowski, Edward 448

Kaplan, E. L. 576

Karlin, John E.

biographical material 589

Reading Rates and the Information Rate of a Human Channel 497-516

KEARNEY WORKS (W. E. Co.) 103-38

Kegelman, T. D. 1126

Kelly, J. L. 512

Kelly, Mervin J.

biographical material 342

Transatlantic Communications—An Historical Resume 1-5

Kelly, R.

biographical material 342

Power-Feed System for the Newfoundland-Nova Scotia Link 277-92

Kelvin, Lord 11, 293

Kip, A. F. 483

KLYSTRON PULLING

interference, interchannel 645-52

Kohman, G. T. 495

Kronacher, Gerald

biographical material 1512

Design, Performance and Application of the Vernier Resolver 1487-1500

L

L-TYPE CARRIER *See* Carrier

LABORATORIES *See* Bell Telephone Laboratories

Lamb, Harold A.

biographical material 342

Repeater Production for the North Atlantic Link 103-38

Lawton, C. S. 1126

LAYING *See* Cable Laying

Leach, Priscilla 1126

LEASED-LINE SERVICES

data transmission

transmission aspects 1451-86

network, shortest connection 1389-1401

Lebert, Andrew W. 495

biographical material 343

Cable Design and Manufacture for the Transatlantic Submarine Cable System 189-216, 496

Lee, C. Y. 1387

Leech, W. H.

biographical material 343

Route Selection and Cable Laying for the Transatlantic Cable System 293-326

Letham, D. L. 512

Levenbach, G. J. 1004

Lewis, Herbert A.

biographical material 343

Route Selection and Cable Laying for the Transatlantic Cable System 293-326
System Design for the North Atlantic Link 29-68

- Lewis, J. A. 495
- LIFE EXPECTANCY
- electron tube
 - 175HQ 166, 171-77
 - 6P12 185-86
 - transatlantic telephone cable
 - North Atlantic link 66-67
- LIMNORIA 1096-1127
- Lince, Arthur H.
- biographical material 344
 - Repeater Design for the North Atlantic Link* 69-101
- LINEAR PROGRAMMING
- binomial processes 537-76
- Lloyd, Stuart P. 964
- Binary Block Coding* 517-35
 - biographical material 589
- LOADING
- repeaters, submarine cable 20
- LONDON 7-9; *map* 8
- Looney, D. H. 1337
- LORAC (navigtion system) 1050
- LORAN (navigation system) 1050
- Loss *See* Net Loss; Transmission Loss
- Lovell, G. H.
- biographical material 344
 - System Design for the North Atlantic Link* 29-68
- Lozier, J. C. 1499
- Lutchko, F. R. 1004
- Lutz, Mary 512
- Lynch, A. C. 495
- M**
- McCall, D. W.
- biographical material 589
 - Determination of Pressure Coefficients of Capacitance for Certain Geometries* 485-95
- McClure, B. T. 768
- McMillan, B. 698
- McNally, J. O.
- biographical material 344
 - Electron Tubes for the Transatlantic Cable System* 163-88
- MAGDALENA RIVER
- turbidity currents 1089-90
- MAGNETIC WIRE *See* Wire
- MAINTENANCE
- carrier, P1 375
 - Newfoundland-Nova Scotia link 235-41
 - North Atlantic link 55-57
 - transatlantic telephone cable 21-23, 55-57, 235-41
- MARINE BORER(s)
- test sites 1115-21
 - transatlantic telephone cable 194
- MARINE NAVIGATION *See* Navigation
- MARITIME PROVINCES OF CANADA 9
- Measurement of Dielectric and Magnetic Properties of Ferromagnetic Materials at Microwave Frequencies* (W. von Aulock, J. H. Rowen) 427-48
- MEMORY ARRAYS
- twistor 1319-40
- Mertz, Pierre
- biographical material 1512
 - Transmission Aspects of Data Transmission Service Using Private Line Voice Telephone Channels* 1451-86
- Meszaros, George W.
- biographical material 345
 - Power Feed Equipment for the North Atlantic Link* 139-62
- METERING
- current, transatlantic telephone cable 151-58
- METHYL-METHACRYLATE *See* Plexiglass
- Metson, G. H.
- biographical material 345
 - Electron Tubes for the Transatlantic Cable System* 163-88
- MICA CAPACITORS *See* Capacitor
- Michaels, S. E. 512
- MICROWAVE(s)
- feed, polarization, dual 391-408
 - paramagnetic resonance techniques 449-84
- MICROWAVE RELAY SYSTEMS *See* Radio Relay Systems
- MID-ATLANTIC RIDGE 1066-68
- MILITARY COMMUNICATIONS
- Nike, electrohydraulic transducer 711-36
 - servomechanisms, hydraulic 736
- Miller, S. E. 405

- MINIATURIZATION, carrier, P1 370
 MISALIGNMENT, transatlantic cable,
 North Atlantic link 42-46
 Mitchel, Duncan M. 754
 Mitchell, Doren
 biographical material 1512
 Transmission Aspects of Data Trans-
 mission Service Using Private Line
 Voice Telephone Channels 1451-86
 MODE(s), normal, electric waves, circu-
 lar 1292-1307
 MODULATION
 amplitude, *see* Amplitude Modulation
 frequency, *see* Frequency Modulation
 MODULATION
 transatlantic telephone cable
 North Atlantic link 63
 pulse
 amplitude (PAM) 655-57
 code (PCM) 655-57
 quantizing impairment 656-57
 duration (PDM) 655-57
 position (PPM) 655-57
 Monarch (cable ship) 162, 244, 250,
 303-26; *illus* 305
 cable gear *line schematic* 310
 MONOGRAPHS, recent, of Bell System
 technical papers not published in
 this Journal 335-37, 583-87; 823-27;
 1043-44; 1313-17; 1508-10
 Monro, S. 576
 MONTREAL 7-9; *map* 8
 Morgan, Samuel P.
 biographical material 1318
 Theory of Curved Circular Waveguide
 Containing an Inhomogeneous Di-
 electric 1209-51
 Mottram, Elliott T.
 biographical material 345
 Transatlantic Telephone Cable System—
 Planning and Over-All Performance
 7-27
 Murphy, R. B. 576
 MUTILATION (data transmission) 1342
- N**
- No. 4B TRANSISTOR *See* Transistor
 No. 5B TRANSISTOR *See* Transistor
 No. 6P10 ELECTRON TUBE *See* Electron
 Tube
 No. 6P12 TUBE *See* Electron Tube
 No. 7F TEST SET *See* Test Set
 No. 175HQ TUBE *See* Electron Tube
 No. SP61 TUBE *See* Electron Tube
 NANTUCKET 13
 NAVIGATION
 marine, systems *table* 1050
 NET LOSS
 transatlantic telephone cable 18
 NETWORK
 carrier, P1 369
 shortest connection 1389-1401
 construction principles 1391-94
 U. S. state capitals *illus* 1390
 switching, periodic 1403-28
 Network Containing a Periodically Oper-
 ated Switch Solved by Successive
 Approximations (C. A. Desoer)
 1403-28
 New Carrier System for Rural Service (R.
 C. Boyd, J. D. Howard, L. Pedersen)
 349-90
 New Storage Element Suitable for Large
 Sized Memory Arrays—the Twistor
 (A. H. Bobeck) 1319-40
 NEW YORK CITY *map* 7-9, 11; 8
 NEWFOUNDLAND *map* 294, 296, 300
 NEWFOUNDLAND-NOVA SCOTIA LINK
 attenuation 221-23
 circuits 223
 crosstalk 230, 243
 design 217-44
 electron tubes 163-88, 179-88
 maintenance 235-41
 noise 229, 241-42
 power supply 225-27, 277-92
 repeaters 163-88, 245-76
 route selection 317-20
 terminals 227-29
 transmission loss 229, 241
 transmission objectives 217
 Niagara (cable ship) 303
 NIKE
 roll servo
 purpose 711
 simplified schematic 712
 transducer, electrohydraulic, J-7
 711-36

NOISE

- electron beam 831-78
- electron tube, traveling-wave 831
- error correction codes 1341-88
- Newfoundland-Nova Scotia Link 229, 241-42
- transatlantic telephone cable
 - North Atlantic link 39-42, 62-63
- radio transmission 623-25
- Noise Spectrum of Electron Beam in Longitudinal Magnetic Field: The Growing Noise Phenomenon; The UHF Noise Spectrum* (W. W. Rigrod) 831-78
- Non-Binary Error Correction Codes* (W. Ulrich) 1341-88
- NONLINEAR ATTENUATION *See* Attenuation
- Normal Mode Bends for Circular Electric Waves* (H.-G. Unger) 1292-1307
- NORTH AMERICA *map* 294, 296
- NORTH ATLANTIC LINK
 - bandwidth 34-35
 - cable current
 - regulation 143-45; *simplified schematic* 143-45
 - channels 38
 - description 29-31
 - design 29-68
 - electron tubes 163-78
 - equalization
 - shore 46-49; *block schematics*
 - inaccessibility 34
 - integrity 33
 - maintenance 55-57
 - misalignment 42-46
 - modulation 63
 - noise 39-42, 62-63
 - performance 59-65
 - power feed, *see* Power Supply
 - repeaters, *see* Repeater
 - schematic diagram* 30
 - signal-to-noise design 35
 - spares 65-66
 - terminals 52-55
 - testing 55-59
- NORTH ATLANTIC OCEAN
 - basins 1064-65
 - bottom 1072-74
 - temperature 1077-86
 - continental shelf 1063-64
 - earthquake epicenters *map* 1087
 - telegraph cables *map* 294
 - topography 1061-70; *illus*

NORTH SEA 3

NORTHERN ELECTRIC COMPANY, LTD.
57, 244

Norton, E. L. 736

NOVA SCOTIA *map* 294, 296

Noyce, R. N. 1004

NUMBER 4B TRANSISTOR *See* Transistor

NUMBER 5B TRANSISTOR *See* Transistor

NUMBER 6P10 ELECTRON GUBE *See*
Electron Tube

NUMBER 6P12 TUBE *See* Electron Tube

NUMBER 7F TEST SET *See* Test Set

NUMBER 175HQ TUBE *See* Electron
Tube

NUMBER SP61 TUBE *See* Electron Tube

O

175HQ TUBE *See* Electron Tube

OBAN, SCOTLAND 2, 9, 29, 49, 57, 140,
145, 147, 150, 164-65, 217, 219, 221,
246, 248, 293, 301, 317-18, 323; *map*
8, 218, 300

OCEAN BOTTOM

catastrophic changes 1086-93

knowledge, present 1070-74

sediment 1071

study 1048

evaluation 1056-1057

methods 1065-70; *illus*

presentation 1057-61

topography 1049-65

turbidity currents 1089-93

OCEAN CABLE *See* Submarine Cable

*Oceanographic Information for Engineer-
ing Submarine Cable Systems* (C. H.
Elmendorf, B. C. Heezen) 1047-93

OCEANOGRAPHY, defined 1047-48

OFFICE *See* Central office

O'Neil, H. T. 1033

OPERATING COMPANIES

carrier, P1 350

ORDNANCE SURVEY OF GREAT BRITAIN
244

P

P1 CARRIER *See* Carrier
 PAM *See* Modulation: pulse; amplitude
 PCM *See* Modulation: pulse: code
 PDM *See* Modulation: pulse: duration
 PPM *See* Modulation: pulse: position
 PARALLEL PLATE CAPACITORS *See* Capacitor
 PARAMAGNETIC RESONANCE
 absorption 450
 PARAMETER(S)
 carrier, P1 351-65
 ferrites 428-30
 PARIS 9
 Pauer, J. J. 754
 Pedersen, Ludwig
 biographical material 589
 New Carrier System for Rural Service 349-90
 PERIODIC SWITCHING *See* Switching
 Perkins, E. H. 390
 PHASE SHIFT
 FM signal, noise modulated
 distortion 879-89
 PHOLADIDAE 1096-1127
 Pierce, John R.
 biographical material 590
 Reading Rates and the Information Rate of a Human Channel 497-516
 PIERCE-TYPE ELECTRON GUN *See* Electron Gun
 PLASTICS
 marine conditions 1095-1127
 PLEXIGLASS
 properties 115
 repeaters, flexible, North Atlantic link 115-16
 POISSON PATTERNS, coincidences 1005-33
 Pollak, H. O.
 biographical material 1045
 Coincidences in Poisson Patterns 1005-33
 POLYETHYLENE
 BOD test 1099
 submarine cable 189-93, 197, 199, 205
 POLYVINYL CHLORIDE, BOD test 1099
 POPULATIONS, binomial *See* Binomial Processes

Portis, A. M. 483
 PORTLAND, MAINE *map* 8
 POST OFFICE *See* British Post Office
 POWER, dc, reliability 140
 POWER-FEED *See* Power Supply
Power Feed Equipment for the North Atlantic Link (G. W. Meszaros, H. H. Spencer) 139-62
Power-Feed System for the Newfoundland-Nova Scotia Link (R. Kelly, J. F. P. Thomas) 277-92
 POWER SUPPLY
 carrier, P1 376-80
 Newfoundland-Nova Scotia link 225-27, 277, 292
 transatlantic telephone cable
 crosstalk 161
 North Atlantic link 49-52, 139-62;
 schematic diagram 51
 equipment design 158-62
 standby sources 145-51
 Prim, R. C.
 biographical material 1512
 Shortest Connection Networks and Some Generalizations 1389-1401
 Prince, M. B.
 biographical material 1045
 High-Voltage Conductivity-Modulated Silicon Rectifier 975-1004
 PRINTED CIRCUITRY
 carrier, P1 371
 PRIVATE LINE SERVICES *See* Leased-Line Services
 PROBABILITIES, binomial processes 537-76
 PROCESSES *See* Binomial Processes
 PROGRAMMING *See* Linear Programming
 PROPAGATION *See* Transmission
 PROSE, information rate 501-04
 PULSE MODULATION, quantized 655-57

Q

QUALITY
 Bell System 103
 Western Electric 103
 QUANTIZED SIGNAL *See* Signal
 QUARTZ CRYSTAL *See* Crystal
 QUEBEC (city) *map* 8

R

- Radio Propagation Fundamentals* (K. Bullington) 593-626
- RADIO RELAY SYSTEMS
- TD-2
- distortion 1429-50
 - equalizer, transversal 1429-50
 - repeaters, equalizer, transversal 1429-50
- RADIO TELEPHONE, transatlantic 5
- RADIO TRANSMISSION LOSS *See* Transmission Loss
- RAIN, and radio transmission 602
- Radley, Sir Gordon
- biographical material 345
 - Transatlantic Communications—An Historical Resume* 1-5
- RAYLEIGH DISTRIBUTION 600-01, 624
- Reading Rates and the Information Rate of a Human Channel* (J. E. Karlin, J. R. Pierce) 497-516
- RECTANGULAR WAVEGUIDE *See* Waveguide
- RECTIFIER
- silicon, conductivity-modulated
 - high-voltage 975-1004
 - solid state
 - voltage 975
- REED-MULLER CODES 1341
- Reflection Theory for Propagation beyond the Horizon* (A. B. Crawford, H. T. Friis, D. C. Hogg) 627-44
- REGENERATIVE REPEATER *See* Repeater
- REGULATOR
- current, transatlantic telephone cable 143-45
- RELAY SYSTEMS *See* Radio Relay Systems
- RELIABILITY
- electron tube, 175HQ 178
- REPEATER
- carrier, P1 362-65, 373-75
 - flexible
 - North Atlantic link 69-138
 - capacitors 125-26
 - circuit 71
 - components 81-88
 - container 90-94
 - coupling networks 73-76
 - design 69-101
 - equalization 99-100
 - feedback loop 78-79
 - gain formula 72-73
 - gas diode tube 88-90; *illus* 89
 - inspection 131-38
 - manufacture 103-38
 - assembly 116-20
 - brazing 116-20
 - clothing, special 109
 - dust count 110
 - quartz crystals 120-23
 - mechanical design 79-81
 - packing 130
 - performance 96-98
 - power feed 139-62
 - production 110-14
 - raw materials 114-16
 - schematic diagram 70
 - seals 90-94, 123-25; *illus* 118
 - shipping 130
 - subcontracted operations 107-08
 - testing 77-78, 94-96
 - Newfoundland-Nova Scotia link 245-76
 - regenerative, self-timing 891-937
 - reliability 245
 - submarine cable
 - British Post Office 12
 - loading 20
 - TD-2 radio relay system
 - equalizer, transversal 1429-50
 - transatlantic telephone cable
 - armoring 58
 - efficiency 2-3
 - electron tubes 2-4
 - specifications 2
 - Repeater Design for the Newfoundland-Nova Scotia Link* (R. A. Brockbank, D. C. Walker, V. G. Welsby) 245-76
 - Repeater Design for the North Atlantic Link* (F. J. Braga, T. F. Gleichmann, A. H. Lince, M. C. Wooley) 69-101
 - Repeater Production for the North Atlantic Link* (W. W. Heffner, H. A. Lamb) 103-38
- RESIN(s)
- casting 1095-1127
 - BOD test 1099
 - marine conditions 1095-1127

Resistance of Organic Materials and Cable Structures to Marine Biological Attack (L. R. Snoko) 1095-1127

RESISTOR

repeater, flexible

North Atlantic link 126-27

RESOLVER *See* Synchro Resolver; Vernier Resolver

RESONANCE 450

See also Paramagnetic Resonance

Rice, Stephen O. 698

biographical material 829, 1046

Distortion Produced in a Noise Modulated FM Signal by Non-Linear Attenuation and Phase Shift 879-89

Interchannel Interference due to Klystron Pulling 645-52

Richards, A. P. 1126

Rigrod, W. W.

biographical material 1046

Noise Spectrum of Electron Beam in Longitudinal Magnetic Field: The Growing Noise Phenomenon; The UHF Noise Spectrum 831-78

Riordan, J. 964-65

RINGING

carrier, P1 359-61

Rose, A. C. 1387

ROUND WAVEGUIDE *See* Waveguide: circular

Route Selection and Cable Laying for the Transatlantic Cable System (J. S. Jack, W. H. Leech, H. A. Lewis) 293-326

Rowen, John H.

biographical material 590

Measurement of Dielectric and Magnetic Properties of Ferromagnetic Materials at Microwave Frequencies 427-48

RURAL TELEPHONE SERVICE

carrier, P1 349-90

S

6P10 ELECTRON TUBE *See* Electron Tube

6P12 ELECTRON TUBE *See* Electron Tube

6P12 TUBE *See* Electron Tube

7F TEST SET *See* Test Set

SP61 TUBE *See* Electron Tube

Schaefer, J. W.

biographical material 830

Electrically Operated Hydraulic Control Valve 711-36

SCOTLAND map 294, 296

SEA *See* Ocean Bottom

SEAL(S), repeater, flexible, North Atlantic link 90-94, 123-25; *illus* 118

SEASONS, transmission, radio, beyond the horizon 640-43

SEDIMENT, ocean bottom 1071

Seidel, Harold

biographical material 590

Character of Waveguide Modes in Gyromagnetic Media 409-26

Selecting the Best One of Several Binomial Populations (Marilyn J. Huyett, M. Sobel) 537-76

Self-Timing Regenerative Repeaters (E. D. Sunde) 891-937

SEMICONDUCTOR(S), SEMICONDUCTING MATERIALS *See* Ferrite

Sensitivity Considerations in Microwave Paramagnetic Resonance Absorption Techniques (G. Feher) 449-84

SERPENTINE WAVEGUIDE *See* Waveguide: circular

SERVICE *See* Maintenance

SERVO SYSTEMS

electrohydraulic, Nike 711-36

hydraulic, military communications 736

power supply, transatlantic telephone cable 155-58

vernier resolver *illus* 1497

SHORAN (navigation system) 1050

SHORT CIRCUIT

cables, Poisson patterns 1005-33

Shortest Connection Networks and Some Generalizations (R. C. Prim) 1389-1401

SIGNAL(S), SIGNALING

binary, data transmission 1451-86

carrier, P1 359

companding, instantaneous 653-709

FM, noise modulated distortion 879-89

quantized, companding, error 665-76
instantaneous 653-709

spectrum 663

transatlantic telephone cable 20-21

- SIMPLEX WIRE AND CABLE COMPANY 196-97
- Silsbee, R. H. 483
- Silverman, S. J. 1004
- Simonick, V. F. 736
- Slepian D. 512
- Slichter, C. P. 483
- Smith, Bernard
biographical material 830
Instantaneous Companding of Quantized Signals 653-709
- Smith, D. H. 390
- Smith, James L.
Activation of Electrical Contacts by Organic Vapors 769-812
biographical material 830
- Snoke, Lloyd R.
biographical material 1318
Resistance of Organic Materials and Cable Structures to Marine Biological Attack 1095-1127
- SNOW and radio transmission 602
- Sobel, Milton
biographical material 590
Selecting the Best One of Several Binomial Populations 537-76
- SO FAR (navigational system) 1050
- SOLENOID electrohydraulic, J-7 711-36
- SOUNDING, echo 1051-52, 1055
- SOUTHERN UNITED TELEPHONE CO., LTD. 244
- SPARE PARTS, North Atlantic link 65-66
- Spencer, H. H.
biographical material 346
Power Feed Equipment for the North Atlantic Link 139-62
- SPRUCE LAKE, NEW BRUNSWICK 11
- ST. JOHN, NEW BRUNSWICK map 8
- STANDARD TELEPHONES AND CABLES, LTD. 244, 274, 292
- STATES (UNITED STATES), capitals, network, shortest connection 1389-1401; *illus* 1390
- STATISTICAL METHODS
telephone exchange model 939-64
traffic fluctuations 965-73
- STORAGE BATTERY, as power source 140
- STORAGE *See* Information Storage
Strength Requirements for Round Conduit (G. F. Weissmann) 737-54
- SUBMARINE CABLE(s)
background experience 11-15
British Post Office systems 12-13
capacitance 485
electron tubes 3
marine organism attack 1095-1127
North Atlantic link 33
oceanographic information 1047-93
recovery
dynamics 1129-1207
kinematics 1129-1207
research 5
stresses 1
transatlantic telephone, *see* Transatlantic Telephone Cable
transistors 3-4
United States 13-14
- SUBMARINE CABLES, LTD. 196-97, 274
- Sufficient Set of Statistics for a Telephone Exchange Model* (V. E. Beneš) 939-64
- Sunde, Erling D. 889
biographical material 1046
Self-Timing Regenerative Repeaters 891-937
- SWITCHING periodic, network 1403-28
- SWITCHING SYSTEMS
electron tubes, gas discharge 755-68
- SWITCHING TIME twistor 1328
- SYNCHRO RESOLVER
accuracy 1487-88
See also Vernier Resolver
- SYDNEY MINES, NOVA SCOTIA 11, 29, 164-65, 219, 221, 232, 246, 248, 318, 321, 323; map 8, 218, 300
System Design for the Newfoundland-Nova Scotia Link (J. F. Bampton, R. J. Halsey) 217-44
System Design for the North Atlantic Link (J. M. Fraser, H. A. Lewis, G. H. Lovell, R. S. Tucker) 29-68
- T**
- TD-2 RADIO RELAY SYSTEM *See* Radio Relay Systems
- TACAN (navigation system) 1050
- TECHNICAL PAPERS, Bell System, not published in this Journal 327-34, 577-82, 813-22, 1035-42, 1308-13, 1501-07

TELEGRAPH

North Atlantic routes *map* 294
transatlantic cable 2, 20, 26-27

TELEGRAPH CONSTRUCTION AND MAINTENANCE CO., LTD. 308

TELEGRAPHIST'S EQUATIONS *See* Generalized Telegraphist's Equations

TELEPHONE EXCHANGE

model, statistics 939-64

TEMPERATURE

North Atlantic Ocean 1077-86
ocean bottom 1075-86

TERMINAL(S)

carrier, P1 354, 373-75
network, shortest connection
1389-1401
Newfoundland-Nova Scotia link
227-29
transatlantic telephone cable
North Atlantic link 52-55

TERRENCEVILLE, NEWFOUNDLAND 11,
219, 232, 246, 248, 293, 295, 318-19,
322, 324; *map* 218, 300

TEST(S), TESTING

biochemical oxygen demand 1098-1114
carrier, P1 375
electron tube, 6P12 186-88
conduit, thin-walled 737-54
repeater, flexible
North Atlantic link 77-78, 94-96

TEST SET

7F 389-90; *illus* 390
conduit, thin-walled *illus* 739

Tharp, Marie 1093

Theory of Curved Circular Waveguide Containing an Inhomogeneous Dielectric
(S. P. Morgan) 1209-51

Thomas, J. F. P.

biographical material 346
Power-Feed System for the Newfoundland-Nova Scotia Link 277-92

TIME OF DAY, and radio transmission
noise 624

Title, R. S. 1337

TONAWANDA PLANT (W. E. Co.) 107

Townsend, Mark A. 88-90

biographical material 830
Cold Cathode Gas Tubes for Telephone Switching Systems 755-68

TRAFFIC

demand 941
engineering 940
fluctuations 965-73
measurement 939-64

Transatlantic Communications—An Historical Resume (M. J. Kelly, Sir G. Radley) 1-5

TRANSATLANTIC RADIO TELEPHONE 5

TRANSATLANTIC TELEPHONE CABLE

cable, *see* Submarine Cable
crosstalk 19-20
echo 21
facilities *block diagram* 10
frequency characteristics 18, 24-26
maintenance 21-23, 55-57, 235-41
map 8, 302
net loss 18
noise 19-20
operating services 21-23
performance 24-27
profile 304
repeaters, *see* Repeater
route selection 293-326
service objectives 16
signaling objectives 20-21
submarine cable, *see* Submarine Cable
system planning 15-24
telegraph facilities 2, 20, 26-27
telephone circuits 9
temperature profile 1083
transmission objectives 16-18

Transatlantic Telephone Cable System—Planning and Over-All Performance
(J. W. Emling, R. G. Griffith, R. J. Halsey, E. T. Mottram) 7-27

TRANSDUCER

coupled-wave, problems 391
electrohydraulic, J-7 711-36
illus 717, 718
actuating mechanism 719-24
cutaway section 716
description 715-19
exploded view 717
hydraulic characteristics 724-34
internal view 720
ports *illus* 714

See also Vernier Resolver

TRANSFORMER, carrier, P1 368-69

TRANSISTOR

- 4B, carrier, P1 355-56
- 4C, carrier, P1 355-56
- carrier, P1 349-90
- submarine cable prospects 3-4
- twistor memory arrays 1333-36

TRANSMISSION

- carrier, P1 351-55
- information, *see* Information Rate
- radio
 - beyond the horizon
 - antenna size 639-40
 - experimental data 608-11
 - reflection theory 627-44
 - seasons 640-43
 - buildings 613-14
 - fog 602
 - fundamentals 593-626
 - ground wave 614-18
 - ionospheric 618-23
 - noise levels 623-24
 - rain 602
 - snow 602
 - trees 613-14
- transatlantic telephone cable 16-18
- Transmission Aspects of Data Transmission Service Using Private Line Voice Telephone Channels* (P. Mertz, D. Mitchell) 1451-86

TRANSMISSION LOSS

- Newfoundland-Nova Scotia Link 229, 241
- radio 593-97
 - earth, plane *diagrams* 598
 - line of sight 596-602

TRANSOCEANIC CABLE *See* Submarine CableTRANSVERSE EQUALIZER *See* EqualizerTRAVELING-WAVE TUBE *See* Electron TubeTREES, and radio transmission 613-14
Tretola, A.R. 1004

TRUNK(S), TRUNKING

- carriers 350
- defined 941-42

TUBE *See* Electron TubeTUBING *See* Copper Tubing

Tucker, Rexford S.

- biographical material 346

System Design for the North Atlantic Link 29-68

Tukey, J. W. 576, 964

TURBIDITY CURRENTS 1089-93

TWISTOR 1319-40

- bits (binary digits) 1336
- switching time 1328
- transistor powering 1333-36

U

USAF *See* United States Air Force

Ulrich, Werner

- biographical material 1513
- Non-Binary Error Correction Codes* 1341-88

Unger, Hans-Georg

- biographical material 1318
- Circular Electric Wave Transmission in a Dielectric-Coated Waveguide* 1253-78
- Circular Electric Wave Transmission through Serpentine Bends* 1279-92
- Normal Mode Bends for Circular Electric Waves* 1292-1307

UNITED STATES

- submarine cable systems 13-14

UNITED STATES AIR FORCE MISSILE TEST CENTER, submarine cable 190-92, 214

V

Van Uitert, L. G. 448

VAPOR, organic, contacts, electrical

- activation 769-812
- erosion 769-812

Vasko, T. J. 1004

Veazie, Edmund A.

- biographical material 347
- Electron Tubes for the Transatlantic Cable System* 163-88

Veloric, Harold S.

- biographical material 1046
- High-Voltage Conductivity-Modulated Silicon Rectifier* 975-1004

VERNIER RESOLVER

- applications 1487-1500
- design 1487-1500
- output 1489
- performance 1487-1500

VERNIER RESOLVER (*Cont.*)

- rotor lamination *illus* 1492
- schematic diagram* 1490
- servo system *illus* 1497
- stator lamination *illus* 1491

VOCABULARY SIZE

- information rate 499

VOCODER channel capacity 497

VOLTAGE rectifier, solid state 975

Volz, A. H. 1449

W

Wakai, T. W. 1499

Walker, D. C.

- biographical material 347
- Repeater Design for the Newfoundland-Nova Scotia Link* 245-76

WAR WORK *See* Military Commun.

Watling, R. G. 754

WASHINGTON, D. C.

- state capitals, shortest connection network 1389-1401; *illus* 1390

WAVE

- circular
 - bends, serpentine transmission 1279-92
 - modes, normal, bends 1292-1307
 - waveguide, dielectric-coated transmission 1253-78
- radio, path 600
- See also* Microwave

WAVEGUIDE

- circular
 - bends, modes, normal 1292-1307
 - birefringence, effect 409-26
 - curved, dielectric, inhomogeneous 1209-51
 - modes, in gyromagnetic media 409-26
 - propagation characteristics 409-26
 - serpentine wave, circular, transmission 1279-92
- coupler, *see* Coupler
- coupling
 - attenuation 392
 - coupler, *see* Coupler
- dielectric-coated, wave, circular, transmission 1253-78

ferrite loaded, propagation characteristics 409

rectangular

- birefringence, effect 409-26
- modes, in gyromagnetic media 409-26

propagation characteristics 409-26

round, *see* Waveguide: circular

Weatherington, C. A. 495

Weissmann, Gerd F.

biographical material 830

Strength Requirements for Round Conduit 737-54

Welber, I. 1449

Welsby, V. G.

- biographical material 347
- Repeater Design for the Newfoundland-Nova Scotia Link* 245-76

Wenny, D. H., Jr. 1337

Werner, J. K. 1449

WEST HAVEN, CONNECTICUT map 8

White, A. D. 768

WHITE PLAINS, NEW YORK map 8

Williams, I. V. 754

Winnicky, A. P. 512

WIRE(S)

magnetic, twistor 1319-40

WIRING, printed, *see* Printed Circuitry

Wittenberg, A. M. 768

Wooley, M. C.

biographical material 347

Repeater Design for the North Atlantic Link 69-101

WORDS

- familiarity, and information rate 500
- length, information rate 500

WRIGHT AIR DEVELOPMENT CENTER 1487, 1499

WRIGHTSVILLE BEACH, NORTH CAROLINA, test site 1115-21

Z

Zadeh, L. A. 1387

Zajac, E. E.

biographical material 1318

Dynamics and Kinematics of the Laying and Recovery of Submarine Cable 1129-1207

ZoBell, Claude E. 1126

